

Coexisting Invariants, Dynamical Spin Control, and Exchangeless Braiding

A Hitchhiker's Guide to Topology

Dissertation

zur Erlangung des Doktorgrades an der
Fakultät für Mathematik, Informatik und Naturwissenschaften
Fachbereich Physik
Universität Hamburg

vorgelegt von

Robin Quade

Hamburg, 2025

Gutachter der Dissertation:	Prof. Dr. Michael Potthoff Prof. Dr. Michael Thorwart
Zusammensetzung der Prüfungskommission:	Prof. Dr. Michael Potthoff Prof. Dr. Michael Thorwart Prof. Dr. Daniela Pfannkuche Prof. Dr. Caren Hagner Prof. Dr. Martin Eckstein
Vorsitzende der Prüfungskommission:	Prof. Dr.. Daniela Pfannkuche
Datum der Disputation:	05.11.2025
Vorsitzender des Fach-Promotionsausschusses PHYSIK:	Prof. Dr. Wolfgang J. Parak.
Leiter des Fachbereichs PHYSIK:	Prof. Dr. Markus Drescher
Dekan der Fakultät MIN:	Prof. Dr.-Ing. Norbert Ritter

Acknowledgements

They say it takes a village to raise a child. Something similar could be said about writing a dissertation. This goes out to those who helped me write mine.

To my first supervisor, Michael Potthoff, who has guided me for more than half my academic life and taught me countless lessons – the two most profound being that A) there is no free lunch, and B) one should always try for free lunch. To my second supervisor, Michael Thorwart, for taking the time to talk physics and academia with me, and for soldering through my magnum opus without a word of complaint. To my committee, for making my defence an educational and fun experience.

To my partner, Cora, whose effortless grace and infinite patience carried me through my worst and elevated me at my best. I may never fully understand how she does it. To my best friend and most steadfast co-conspirator through all of uni, Nicolas, because he made me mention him. To the members and associates of our research group – Cassian, Simon, Nicolas, David, Leon, Micha, Christian, Michael, Marvin, Pia, Sarah, Devesh, Angela and Lara – without you, I would probably have been finished by 2024, but what fun would that have been? Lara, the cake was exquisite.

To my parents, who funded my undergraduate studies and made many attempts at taking an interest in topological quantum matter; this is how much I am loved. To my sister, whose dedication as a mother inspires me daily, and to my niece and nephew, who do not (yet) care for physics but have taught me much about its limits even so. To everyone else, from bandmates to old and new friends, for keeping me sane and nodding along to my tales and puns at parties. To my Bucerius crew, Leon and Isabella, for sweetening the final stretch considerably.

To my obsessive tendencies, which, after decades of pestering me for sport, I have somehow tricked into serving me with a fierce tenacity. To evolution – for cats. And finally, to music.

Mischief managed.

Abstract

Modern condensed matter theory increasingly draws on topological ideas to explain phenomena that defy purely local descriptions. While the theoretical foundations of topology in quantum matter are well established, much can still be learnt about how topological properties interact with other aspects of physical systems. This thesis explores the interplay between topology and competing physical effects through a series of quantum and quantum-classical case studies. We cover systems with coexisting topological orders, quantum-classical dynamics, and small-scale quantum information platforms.

The first study addresses the implications of coexisting topological orders in a quantum-classical hybrid system formed by classical impurity spins coupled to a spinful Haldane model on a periodic honeycomb lattice. In this setting, the configuration space of the classical impurity spins constitutes a parameter manifold that enables an extrinsic monopole-like classification complementing the intrinsic Chern classification of the Haldane host. The additional monopole-like topological order explains the emergence of a spectral flow of bound-state energies bridging the host's energy gap as a function of the exchange-coupling strength. The form of this spectral flow, however, is determined by the Chern topology of the host system, which dictates whether or not impurity-bound in-gap states appear in the limit of infinitely large exchange coupling strengths between impurity spins and host. The coexistence of the conventional Chern momentum-space topology and the monopole-like configuration-space topology enables the numerical construction of topological phase diagrams that provide valuable insights into the interrelations between the two topological structures.

The second study demonstrates how the topological edge states of a quantum spin Hall system can be harnessed to control the real-time dynamics of a magnetic impurity. To this end, a classical impurity spin is coupled to one edge of a Kane–Mele model on a finite ribbon segment. The real-time evolution of the resulting quantum-classical hybrid system is obtained by numerically solving the full set of coupled equations of motion for the classical spin and the electronic system. In order to enable the simulation of long-time dynamics, dissipative boundary conditions are imposed on all but the impurity-hosting edge. Spin density injections into the spin-momentum locked Kane–Mele edge modes propagate unidirectionally and with minimal loss, enabling remote manipulations of the impurity spin dynamics. Iterated protocols of spin injection, unidirectional propagation, and scattering off the impurity spin allow the implementation of a complete and reversible spin switching process.

The third study employs superconducting phase rotations in small networks of weakly-linked Kitaev chains to realise logical braiding operations for topological quantum computation without exchanging anyonic Majorana quasiparticles. Braiding protocols are verified by numerically evaluating the non-Abelian Wilczek–Zee phase of the low-energy many-body subspace $\mathcal{H}_0(\phi)$ using the Bertsch–Robledo overlap formula. The parameter space of two weakly linked Kitaev chains divides into two regions characterised by distinct braiding outcomes, i.e. projective σ_x - or σ_z -type braiding operations. A selection of representative phase diagrams reveal that the σ_x and σ_z phases are separated by a continuous crossover. Moreover, this crossover regime is centred on a sharp transition hypersurface derived from a minimal model of the four involved Majorana modes. The study demonstrates the robustness of anyonic properties against finite-size effects and weak couplings between Kitaev chains.

Kurzzusammenfassung

Die moderne Theorie der kondensierten Materie greift zunehmend auf topologische Ideen zurück, um Phänomene zu erklären, die sich rein lokalen Beschreibungen entziehen. Während die theoretischen Grundlagen der topologischen Quantenmaterie gut etabliert sind, kann noch viel darüber gelernt werden, wie topologische Eigenschaften mit anderen Aspekten physikalischer Systeme interagieren. In dieser Dissertation wird das Zusammenspiel zwischen Topologie und konkurrierenden physikalischen Effekten anhand einer Reihe von quantenmechanischen und quantenklassischen Fallstudien untersucht. Diese umfassen Systeme mit koexistierenden topologischen Ordnungen, quantenklassische Dynamik und kleine Quanteninformationsplattformen.

Die erste Studie befasst sich mit den Auswirkungen koexistierender topologischer Ordnungen in einem quantenklassischen Hybridsystem, in dem klassische Störstellenspins an ein Haldane-Modell von Elektronen mit Spin-1/2 auf einem periodischen Honigwabengitter gekoppelt sind. In diesem Rahmen stellt der Konfigurationsraum der klassischen Störstellenspins eine Parametermannigfaltigkeit dar, die eine extrinsische monopolartige Klassifizierung ermöglicht, welche die intrinsische Chern-Klassifizierung des Haldane Trägersystems ergänzt. Die zusätzliche monopolartige topologische Ordnung erklärt das Auftreten eines spektralen Flusses von Energien gebundener Zustände, welche die Energielücke des Haldane-Modells als Funktion der Austauschkopplungsstärke überbrücken. Die Form dieses spektralen Flusses wird derweil durch die Chern-Topologie des Trägersystems bestimmt, die festlegt, ob im Grenzfall unendlich großer Austauschkopplungsstärken zwischen den Störstellenspins und dem Trägersystem gebundene In-Gap-Zustände auftreten oder nicht. Die Koexistenz der konventionellen Chern Impulsraumtopologie und der monopolartigen Konfigurationsraumtopologie ermöglicht die numerische Konstruktion topologischer Phasendiagrammen, die wertvolle Einblicke in die wechselseitigen Beziehungen zwischen den beiden topologischen Strukturen liefern.

Die zweite Studie zeigt, wie die topologischen Randzustände eines Quanten-Spin-Hall-Systems zur Steuerung der Echtzeitdynamik einer magnetischen Störstelle genutzt werden können. Zu diesem Zweck wird ein klassischer Störstellenspin an eine Kante eines Kane–Mele-Modells auf einem endlichen Streifensegment gekoppelt. Die Echtzeitentwicklung des resultierenden quantenklassischen Hybridsystems wird ermittelt, indem der vollständige Satz gekoppelter Bewegungsgleichungen des klassischen Spins und des elektronischen Systems numerisch gelöst wird. Um eine Simulation der Langzeitdynamik zu ermöglichen, werden dissipative Randbedingungen für alle Ränder mit Ausnahme des Randes, der den Störstellenspin trägt, eingeführt. Die feste Kopplung zwischen Spin- und Impulsrichtung in den Kane–Mele-Randmoden führt dazu, dass sich Spindichte-Injektionen unidirektional und nahezu verlustfrei ausbreiten, was eine gezielte Manipulation der Störstellenspindynamik aus der Ferne ermöglicht. Iterierte Protokolle aus Spininjektion, unidirektionaler Ausbreitung und Streuung am Störstellenspin ermöglichen die Implementierung eines vollständigen und reversiblen Spinschaltprozesses.

Die dritte Studie nutzt Rotationen der supraleitenden Phase in kleinen Netzwerken aus schwach verknüpften Kitaev-Ketten, um logische Braidingoperationen für topologisches Quantencomputing zu realisieren, ohne die anyonischen Majorana-Quasiteilchen auszutauschen. Die Braidingprotokolle werden numerisch verifiziert, indem die nicht-abelsche Wilczek–Zee-Phase des niederenergetischen Vielteilchen-Unterraums $\mathcal{H}_0(\phi)$ unter Verwendung der Bertsch–Robledo-Überlappformel ausgewertet wird. Der Parameterraum von zwei schwach verknüpften Kitaev-Ketten teilt sich in zwei Regionen, die durch unterschiedliche Braidingresultate gekennzeichnet sind, nämlich projektive σ_x - oder σ_z -artigen Braidingoperationen. Eine Auswahl repräsentativer Phasendiagramme zeigt, dass die σ_x - und σ_z -Phasen durch einen kontinuierlichen Übergang voneinander getrennt sind. Dieser kontinuierliche Übergang ist darüber hinaus auf einer scharfen Übergangs-Hyperfläche zentriert, welche aus einem Minimalmodell der vier beteiligten Majorana-Moden abgeleitet wird. Die Studie zeigt, dass die anyonischen Eigenschaften selbst bei endlichen Systemgrößen und schwacher Kopplung zwischen den Kitaev-Ketten erhalten bleiben.

Table of Contents

1	Introduction	1
2	Mathematical Background	6
2.1	Topology	6
2.1.1	Topological Spaces	6
2.1.2	Continuity	7
2.1.3	Separation	8
2.1.4	Connectedness	9
2.1.5	Compactness	9
2.1.6	Homeomorphisms	10
2.1.7	(The Hole Truth About) Topological Invariants	11
2.1.8	Homotopy	12
2.1.9	Homology	16
2.1.10	Cohomology	25
2.2	Fibre Bundles	33
2.2.1	Vector Bundles	36
2.2.2	Principal Bundles	37
2.2.3	Pullback Bundles and Classifying Spaces	38
2.2.4	Connections on Fibre Bundles	40
2.2.5	Holonomy	46
2.2.6	Curvature	48
2.3	Characteristic Classes and Numbers	52
2.3.1	The Euler Class	53
2.3.2	The Stiefel–Whitney Classes	55
2.3.3	The Chern Classes	56
2.3.4	The Pontryagin Classes	57
2.3.5	Chern–Weil Theory	58
2.3.6	Chern Numbers and a Glimpse of the Atiyah–Singer Index Theorem	60
3	Time Reversal and Particle Hole Symmetry	63
3.1	Time Reversal Symmetry	64
3.2	Particle Hole Symmetry	66
4	Geometric Phases	68
4.1	The Foucault Pendulum	68
4.2	The Berry Phase	69
4.3	The Wilczek–Zee Phase	74
4.4	Berry, Aharonov and Anandan	76
4.5	Geometric Features of Diabatic Dynamics	79
4.6	Discretised Geometric Phases	80

5	Bogoliubov–de Gennes Theory	83
5.1	Bardeen–Cooper–Schrieffer Theory of Superconductivity	83
5.2	The Bogoliubov–de Gennes Formalism	87
5.3	The Bogoliubov Vacuum	100
5.4	Bogoliubov Overlaps	109
5.5	Majorana Modes in BdG Theories	113
6	Anyons	118
6.1	Anyons and Spin	123
6.2	Topological Quantum Computation with Anyons	123
6.3	Algebraic Theory of Anyons	124
6.3.1	Fibonacci Anyons	127
6.3.2	Ising Anyons	129
6.4	Emergent Anyons in Topological Superconductors	131
6.5	Topological Quantum Computation with Anyons Revisited	138
7	Spectral Impurity Responses in Systems with Coexisting Topological Structures	139
7.1	Multi-Impurity Kondo Model with Classical Spins	140
7.1.1	S -Space Topology of the Multi-Impurity Kondo Model with Classical Spins	141
7.2	The Haldane Model	141
7.2.1	The Graphene Lattice	142
7.2.2	Topology of the Haldane Model	144
7.2.3	Spin Haldane Model with Classical Spin Impurities	146
7.3	Low-Energy Electronic Structure with a Single Impurity Spin	146
7.4	Strong- J Limit: The Magnetic Monopole	148
7.5	Topological Transition at J_{crit}	150
7.6	Relation to k -Space Topology	151
7.7	Two Impurity Spins	154
8	Long-Range Helical Spin Control	159
8.1	The Kane–Mele Model	160
8.1.1	Topology of the Kane–Mele Model	163
8.1.2	Bulk-Boundary Correspondence in Zigzag Ribbons	165
8.2	Real-Time Dynamics in the Kane–Mele s-d-Model	167
8.3	Macroscopic Edge Dynamics in Ribbon Segments with Lindblad Boundaries	169
8.4	Spin Injection and Helical Transport	170
8.5	Helical Transport in the Presence of the Read-Out Spin	173
8.6	Helical Control over a Read-Out Spin	174
8.7	Iterated Injection Protocol: Building a Helical Spin Switch	176
9	Sombrero Berry-Phases in BdG Vacuum Manifolds	180
9.1	Bogoliubov Diagonalisation of Reduced BdG Hamiltonians	180
9.2	The Sombrero Berry-Phase of the BCS Ground State	181
9.3	The Sombrero Berry-Phase of the Kitaev Chain Ground State	183

10	Exchangeless Braiding with Non-Degenerate Anyons	185
10.1	The Kitaev Chain Model	186
10.1.1	Topology of the Kitaev Chain	190
10.2	Computational Methodology	193
10.3	Exchangeless Braiding in One Kitaev Chain	194
10.4	Exchangeless Braiding in Two Kitaev Chains	198
11	Conclusion and Outlook	206
A	Appendix	212
A.1	The Tenfold Way and Real Division Algebras	212
A.2	The Real Schur Decomposition Theorem	215
A.3	Validity of the Thouless Vacuum	217
A.4	Skew-Symmetry of the S -Matrix	219
A.5	The Bloch–Messiah Decomposition	220
A.6	The Relation Between the Product State and the Thouless State	229
A.7	The Robledo Overlap Formula for Product States	231
A.8	TRS, Fourier Transform and Topology of the Haldane Model	233
A.9	The Magnetic Monopole	241
A.10	TRS, Fourier Transform and Topology of the Kane–Mele Model	251
A.11	Diagonalisation and Product Vacuum of BCS and Kitaev Chain Hamiltonians	256
A.12	Technical Aspects of the Kitaev Chain Model	261
B	Bibliography	275
C	List of Publications	288



1 – Introduction

Long before humanity formalised its pursuit of knowledge through empiricism, rationalism, and later the scientific method, it entertained an obsessive fascination with *patterns* – carved into cave walls to capture the archetypical features of predator and prey; shaped into the sounds and symbols of language; recorded in the early calendars charting the celestial cycles of the sun, the moon, and the stars. The human interest in such regularities is hardly surprising: faced with the overwhelming complexity of nature, patterns have always created a sense of order, something to hold on to amid the chaos.

From this viewpoint, one might argue that the development of mathematics as a formal system to describe and abstract patterns was more or less inevitable. Likewise, it is perfectly intuitive that such a system would find application in the modern natural sciences. What is much less intuitive, is the consistency with which mathematics assumes a role that goes far beyond mere utility: shaping not only how we *describe* the world, but how we *discover* it. This remarkable effectiveness inspired Wigner’s now-famous essay, “The Unreasonable Effectiveness of Mathematics in the Natural Sciences”, in which he expressed his astonishment, even disbelief, at how uncannily well mathematics fits the natural world [1]. Wigner illustrated his argument with some of the most-prevalent applications of mathematics in physics at the time: differential calculus in classical mechanics, functional analysis in quantum mechanics, and group theory in quantum electrodynamics. Had he written his essay a couple of decades later, he might have featured another branch of mathematics: the late 20th-century rise of topology – subtle, slippery, and suddenly everywhere in physics – would have fit his case perfectly.

A Minimal Taxonomy of Topology in Physics. In hindsight, topology has been lurking in physics since Dirac’s magnetic monopole in 1931 [2], the Aharonov–Bohm effect in 1959 [3], and nuclear skyrmions in the early 1960s [4, 5]. Yet, it was not until the 1970s and 1980s, with the first explicit mentions of topological solitons in field theory [6–8], the topological characterisation of the quantum Hall effect [9–11], and the introduction of topological quantum field theories [12, 13], that the scientific community recognised “topological effects” as a distinct class of physical phenomena. While the details of such topological effects may vary greatly, they all share a simple unifying theme: they possess a quality that is in some sense *global* and invariant under *continuous* transformations. This could be a *twist* in the Bloch states over reciprocal space, but also a *winding* of a classical field configuration in real space, or an *obstruction* in a gauge bundle over spacetime. Since the already long list of topological phenomena continues to grow at a remarkable pace, devising a comprehensive list of topology-related physical effects is a hopelessly difficult task. For this reason, we content ourselves with assembling a collection of the most prevalent flavours, a *minimal taxonomy*, of topology in physics.

Quantum Matter. Topological quantum matter is generally concerned with the topological properties of quantum states in macroscopic lattice systems of free fermions. The translational invariance in the bulk of such systems endows its quantum states with a (Bloch-)bundle structure, which enables an extensive topological classification in terms of characteristic classes and K -theory [14, 15]. In its simplest form, this classification tells us how the occupied Bloch bands of a band insulator twist as a function of quasi-momentum. The resulting phase of matter is known as a *topological insulator* (TI) and is usually characterised by a single $\mathbb{Z}$, $\mathbb{Z}_2$, or $2\mathbb{Z}$ number known as its topological invariant. Importantly, the topology, that the many-body ground state of a TI inherits from its occupied Bloch bands, can only change when the insulating band gap closes. In this sense, the band gap of a TI protects the topological properties of its ground state. The same line of reasoning can be applied to ground states of Bogoliubov–de Gennes (BdG) superconductors, giving rise to *topological superconductors* (TSCs) whose ground state topology is then protected by the superconducting gap. Another example of topological matter are Weyl semimetals (WSMs). Unlike TIs and TSCs, these arise from gapless bulk band structures where conduction and valence bands touch at isolated Weyl points that carry a topological charge. Depending on whether the topological properties of a given phase of matter depend on the presence of symmetries, it is further distinguished between symmetry-protected and intrinsic topological phases. Notably, the topological classification of Bloch bundles in periodic lattice models can be generalised through non-commutative geometry to encompass more realistic, disordered crystals [16–18].

Many topological phases feature boundary modes, such as edge states of two-dimensional TIs [19–21] or surface Fermi arcs of three-dimensional WSMs [22–24]. These often possess exotic properties and can serve as signatures of the bulk topology. The existence of such boundary modes is typically guaranteed by a principle called bulk-boundary correspondence [14, 25]. More generally, the tenfold way of topological quantum matter with symmetries [14] classifies the bulk-*defect* correspondence of free-fermion topological phases of matter. It is based on the ten Altland-Zirnbauer classes X of fundamental symmetries (time-reversal, particle-hole, chiral), the system dimensions d , and the defect codimensions D [26, 27]. For each of these combinations, the classification scheme specifies whether a generic codimension- D defect in a system of spatial dimension d and symmetry class X supports topological zero modes, and if so, which type of bulk topological order ($\mathbb{Z}$, $\mathbb{Z}_2$, $2\mathbb{Z}$) protects them via bulk-defect-correspondence.

Exemplary models that will be discussed in this thesis include the Kitaev chain, a one-dimensional TSC with particle-hole symmetry and Majorana zero modes, the Haldane model, a two-dimensional TI without symmetries and chiral edge modes, and the Kane–Mele model, a two-dimensional TI with time-reversal symmetry and helical edge modes.

Solitons. The first documented evidence of solitonic behaviour dates back to the early 19th century, when Scottish naval architect John Scott Russell reported a puzzling type of water wave that maintained its shape over surprisingly long distances and even through collisions with other waves [28]. Russell’s fascination with what he called the “solitary wave of translation” was way ahead of its time. In fact, its broader significance for physics was not recognised until about a century later, when American physicist Norman J. Zabusky introduced the term *soliton* to describe stable, localised waves in one-dimensional non-linear dispersive media [29]. Zabusky chose the name as a blend of “solitary” and the particle suffix “-on”, reflecting the wave’s unusual ability to propagate in isolation and without changing its shape; much like a stable particle. Russell’s solitary wave of translation and Zabusky’s one-dimensional solitons have since evolved into the much more general and surprisingly universal concept of topological solitons.

In the spirit of the original phenomena, modern topological solitons refer to stable, particle-like configurations of classical fields $\phi : \mathbb{R}^d \rightarrow E$, which are defined over d -dimensional Euclidean space $\mathbb{R}^d$ and take values in some manifold or linear space E . They occur in many non-linear classical field theories all across physics [30]. More concretely, solitonic field solutions are characterised by a strong spatial localisation and a striking resilience against dispersion, distortion, and annihilation. Remarkably, this stability comes from a *topological* obstruction: a soliton exhibits a non-trivial global winding that distinguishes it from the trivial vacuum of the theory. Since this winding cannot be undone by any continuous deformation, the soliton cannot continuously decay into that vacuum. In more technical terms, the vacuum and solitons belong to different *homotopy* classes. Each homotopy class corresponds to an equivalence class of continuous maps from the field domain $\mathbb{R}^d$ into the codomain E . The most fundamental requirement for the existence of topological solitons is therefore a partial-differential equation that possesses topologically distinct homotopy classes of solutions. In many cases, these classes arise because the topologically trivial field domain $\mathbb{R}^d$ can be given the topologically non-trivial structure of a sphere, either by one-point compactification $\mathbb{R}^d \cup \{\infty\} \simeq \mathbb{S}^d$ or by identifying its asymptotic boundary as $\partial_\infty \mathbb{R}^d \simeq S_\infty^{d-1}$ [30]. This is possible because one is typically interested in fields ϕ that fulfil physical boundary conditions at infinity, i.e. approach a direction-independent constant value $e \in E$ or realise direction-dependent values in a vacuum manifold $\mathcal{V} \subset E$ [30]. In either case, the homotopy classes of compactified or asymptotic fields $\bar{\phi} : \mathbb{S}^d \rightarrow E$ or $\phi_\infty : S_\infty^{d-1} \rightarrow \mathcal{V}$ define elements of so-called homotopy groups $\pi_d(E)$ or $\pi_{d-1}(\mathcal{V})$, providing the topological structure needed to distinguish field solution classes.

The resulting *soliton sectors* of the field theory are thus characterised by an element of the relevant homotopy group. Often, this homotopy group is isomorphic to $\mathbb{Z}$ or $\mathbb{Z}_2$ and the group elements characterising the soliton sectors are called soliton numbers N or, for reasons that will become clear shortly, topological charges. It is worth noting that the asymptotic homotopy classification via $\pi_{d-1}(\mathcal{V})$ requires a non-trivial vacuum manifold $\mathcal{V} \subset E$. Such a manifold corresponds to a continuum of degenerate vacuum states, indicating spontaneous symmetry breaking. For this reason, soliton physics is often associated with spontaneous symmetry breaking.

Due to their great conceptual generality, topological solitons appear in numerous areas of physics. In quantum field theory they manifest as skyrmions of non-linear sigma models [4, 5] and instantons in some gauge theories [7, 30–32].¹ In condensed matter theory, they emerge in the form of magnetic skyrmion textures [33], topological lattice defects like dislocations [34, 35], and vortices in superfluid phases [36–38]. Beyond this, solitons describe monopoles like the magnetic Dirac monopole [2, 30, 39] or the ’t Hooft–Polyakov monopole [6, 40], domain walls in magnetism [41, 42] and optics [43, 44], as well as solitary waves in fluid dynamics [28, 45].

In the present work, solitons appear only implicitly, in the form of half-quantum flux vortices hosting Majorana zero modes in two-dimensional $p_x + ip_y$ superconductors.

Conservation Laws. The standard route to conservation laws in physics is via Noether’s theorem: given a continuous symmetry of an action $\mathcal{S}$, one obtains a Noether current j built from the fields ϕ and their derivatives $\partial_\mu \phi$. This current is conserved on-shell, meaning $\partial_\mu j^\mu = 0$ whenever ϕ satisfy the Euler–Lagrange equations. The Noether charge $q = \int_\Sigma j^0 d^3x$ is then conserved, $dq/dt = 0$, on all constant-time slices Σ . Now, the existence of solitons demonstrates that topology can also conserve certain features of a field theory. Indeed, there exists a second type of physical conservation laws that is closely related the notion of both Noether conservation laws and topological solitons [46]. These *topological* conservation laws are not tied to symmetries of an action, but to the global topology of the solution space of its Euler–Lagrange equations. The conserved topological charge, $Q = \int_{\mathbb{S}^3} j^0 d^3x = 4\pi N$, is essentially determined by the soliton number N mentioned in the previous section. It can be used to define a conserved topological current J satisfying $\partial_\mu J^\mu = 0$. Note that unlike Noether currents and charges, topological currents and charges are invariant under continuous deformations of ϕ , and consequently hold even off-shell [46].

Index Theorems. Index theorems establish deep connections between the analytical and topological properties of certain differential operators $D : \Gamma(E_1) \rightarrow \Gamma(E_2)$, where $E_i \xrightarrow{\pi_i} M$ are vector bundles over a compact manifold M and $\Gamma(E_i)$ denotes the set of smooth sections of E_i . More concretely, they equate an *analytical index*, which measures the net imbalance between the number of solutions and the number of constraints, to a *topological index*, which is defined purely in terms of global topological data [39]. The most prominent index theorems are the Atiyah–Singer index theorem [47] and its extension to manifolds with boundary, the Atiyah–Patodi–Singer index theorem [48–50]. Some topological phenomena in physics emerge when an index theorem ensures the existence of zero-energy solutions to a given (often Dirac-type) differential equation, thereby giving them a topological character.

For instance, the bulk topological order of a condensed matter system can give rise to defect-bound zero modes through the aforementioned bulk–defect correspondence [14]. Mathematically, this correspondence can be formalised using index theorems like the Atiyah–Singer or Atiyah–Patodi–Singer theorems, which relate the existence of such zero modes to global topological invariants of the bulk system [14, 25]. The chiral edge modes of the quantum Hall effect, the helical edge modes of the quantum spin Hall effect, and the Majorana zero modes of topological superconductors are prominent examples of this.

In a strikingly similar fashion, *oceanic* Kelvin waves can be understood as topological edge modes. In Ref. [51], Delplace *et al.* draw a surprising parallel between Kelvin waves and the chiral edge modes of the integer quantum Hall effect: in a rotating shallow-water system, the Coriolis force acts like a magnetic field, breaking time-reversal symmetry and enabling a Chern classification of the bulk Poincaré wave modes. Perhaps counterintuitively, this shallow-water model provides a good approximation of oceanic fluid dynamics. The reason for this is that the ocean is a fluid layer whose horizontal extent vastly exceeds its vertical height. As a result, one can effectively treat it as a two-dimensional “Chern fluid” on a sphere. From this perspective, the equator represents an interface between regions of opposite Coriolis forces and Chern numbers, hosting chiral edge modes, i.e. equatorial Kelvin waves, which propagate eastward with opposite chiralities in each hemisphere. The opposite chiralities of the northern and southern hemispheres edge modes also manifests along topographic boundaries such as coastlines, where coastal edge modes propagate in counterclockwise (clockwise) direction in the northern (southern) hemisphere. A particularly elegant way of making the topological nature of the Kelvin waves manifest is to reduce the shallow water equations to a relativistic Maxwell–Chern–Simons gauge theory, as was done in Ref. [52].

¹Skyrmions are named after Tony Skyrme, who first proposed them as a model for the nucleon in 1961 [4]. Instantons are named for being localised in both space and time, effectively describing an event confined to a point in space and an *instant* in time.

Finally, index theorems form the mathematical basis of quantum anomalies, i.e. the anomalous non-conservation of classically conserved currents in quantum field theories [53–56]. A well-known example is the Adler–Bell–Jackiw chiral anomaly in quantum electrodynamics, which explains the anomalous decay rate of the neutral pion via the non-conservation of the axial or chiral current [57, 58]. Conceptually, this non-conservation is a result of tunneling between topologically distinct vacua [59], a process generated by instantons of the gauge field [59].

Topology: Use, Abuse, and Beyond. The above exposition places us in the midst of a veritable zoo of topological effects in physics, raising the question of what we can do with these. While the topological classification of physical phenomena is a considerable achievement in its own right, we are, from a physicists’ perspective, compelled to look for concrete physical use cases and implications. Of course, the preceding taxonomy of topological phenomena already included several useful implications of this kind: bulk-boundary correspondence in topological quantum matter, quantum anomalies in gauge theories, and soliton solutions in fluid dynamics clearly constitute tangible physical effects of specific practical or theoretical value.

Yet, the promises of topology do not come without peril either: so seductive is the notion of topological order that it is in constant danger of devolving into a fashion label worn more for flair than substance. In this light, it is also the task of physicists to remain critical and, if necessary, to resist the allure of topology in order to preserve the grounds for meaningful discourse. Not every global property is topological, and not every defect mode owes its existence to bulk-defect correspondence.

Here, we set out to explore a variety of concrete use cases for topology that build upon and go beyond the aforementioned foundational applications. In doing so, we take particular care to continually address and scrutinise the role of topology. Specifically, we examine three distinct frontiers of topology:

- * the *Topology and Topology* frontier, which characterises systems with coexisting or competing topological structures,
- * the *Topology and Dynamics* frontier, which addresses the influence of topological defect modes on the dynamics of local impurities, and
- * the *Topology and Quantum Computation* frontier, which concerns the persistence and operability of topological qubits in small-scale quantum systems.

Organisation of this Thesis. As a foundation, Chapter Two provides a compact but thorough mathematical review of the topological, algebraic and geometrical concepts underlying the phenomena studied throughout the thesis. Aimed at readers with a background in physics, most of the material is reviewed from the ground up. It covers basic notions from general and algebraic topology, the theory of fibre bundles, and characteristic classes, culminating in an excursion to the Atiyah–Singer index theorem. The mathematical groundwork discussed in this chapter enables a rigorous exploration of the topological structure of quantum matter, while maintaining a clear focus on the physical content throughout the remainder of the thesis.

The mathematical review is followed by a series of chapters presenting more specialised concepts and methods from theoretical physics. As a start, Chapter Three provides a brief guide to time-reversal and particle-hole symmetry, facilitating the discussion of time-reversal protected quantum spin Hall systems and particle-hole symmetric BdG superconductors. Following that, Chapter Four gives an overview of Abelian and non-Abelian geometric phases. These supply the technical foundation for the later analysis of braiding between anyonic Majorana zero modes. Chapter Five offers an in-depth discussion of the BdG formalism, laying the groundwork for the numerical implementation and analysis of topological superconductivity in the Kitaev chain. Finally, Chapter Six gives an introduction to the theory of algebraic anyons, preparing a rigorous description of Majorana zero mode statistics.

Chapter Seven is the first of three chapters centred on numerical results. It is dedicated to the *Topology and Topology* frontier, presenting a numerical study concerning the interplay and conceptual value of coexisting topological structures. This is done by exchange-coupling a small number of R classical impurity spins to a two-dimensional topological Chern insulator with periodic boundary conditions. The configuration space of the R impurity spins forms a closed $2R$ -dimensional parameter manifold and enables a topological classification in terms of what we call the R -th *spin*-Chern number $Ch_R^{(S)}$, complementing the conventional bulk classification of the translationally invariant host system by means of the first k -Chern number $C_1^{(k)}$. The results of this chapter are published in [RQ1].

The subsequent Chapter Eight explores the frontier of *Topology and Dynamics*. Based on a quantum-classical hybrid system in which a classical impurity spin is coupled to the edge of a two-dimensional quantum spin Hall insulator, we investigate how helical boundary modes affect the impurity spin dynamics. A macroscopic half-space geometry is simulated by exchange-coupling the classical impurity spin(s) to one edge of a finite ribbon segment and applying absorbing boundary conditions along the remaining edges. This makes it possible to analyse the long-time dynamics of both the electronic host system and the classical impurity spin by solving the coupled equations of motion of the quantum-classical hybrid system. The results of this chapter are published in [RQ2].

Chapter Nine and Chapter Ten are concerned with the *Topology and Quantum Computation* frontier. Methodologically, both chapters focus on a numerical analysis of “Sombrero” Wilczek–Zee phases characterising the geometric evolution of superconducting many-body states under full rotations of the superconducting phase. Chapter Nine addresses the Abelian Sombrero Wilczek–Zee phases, i.e. the Sombrero *Berry* phases, associated with the unique ground states of conventional *s*-wave and topological *p*-wave superconductors. Chapter Ten builds on this by examining the non-Abelian Sombrero Wilczek–Zee phases associated with the almost degenerate ground states of a topological *p*-wave superconductor hosting Majorana zero modes. This analysis allows us to decode the anyonic statistics of the Majorana zero modes. The results of this chapter are published in [RQ3].

A full list of the publications supporting this work, together with a detailed explanation of the author’s contributions, can be found in the Declaration of Publications.

2 – Mathematical Background

In an insight far ahead of its time, Galileo famously described mathematics as the language in which the universe is written [60]. He claimed that without mathematics, one was left to wander about in a dark labyrinth. Even if Galileo’s artful prose may not have stood the test of time,¹ it is safe to say that the ideas he put forward have. In modern theoretical physics, mathematics has long advanced from a mere framework for rigour to a guiding light that reveals new patterns and sparks unexpected insights. One area for which this is particularly true is that of topological phenomena in physics.

This chapter introduces the mathematical concepts underlying the physical phenomena considered in the following thesis. Specifically, we cover the fundamental definitions and notions from topology, fibre bundles, and characteristic classes. The presentation closely follows Refs. [15, 39, 61].

2.1 Topology

Mathematics is often associated with intricate detail and sophistication. In a sense, topology goes against this notion: it adopts a deliberately blurred perspective, asking which structural details can be ignored to get to the most fundamental properties. This principle is often visualised as rubber-sheet geometry, where objects can be stretched and deformed like rubber, but can never be torn or glued. Under these conditions, detailed geometric features like the local curvature or volume lose their significance, leaving only topological properties such as the famous number of holes in an object.

In this chapter we give an introduction to topology, covering concepts from both general and algebraic topology. This part is largely based on Refs. [39] and [61].

2.1.1 Topological Spaces

The name *topology* is a compound term made up of the Greek *topos* meaning *place* and *logos* meaning “reason”. Accordingly, topology is the study of location or proximity. Indeed, all of the fundamental concepts of general topology, like continuity and connectedness, revolve around a notion of proximity. The most essential task of general topology is to provide a purely set-theoretical foundation for it [39].

Definition 2.1.1. Let $I, J \subseteq \mathbb{N}$ denote suitable index sets. Let X be some set. A **topology** $\mathcal{T}$ on X is a family $\mathcal{T} = \{U_i \subseteq X \mid i \in I\}$ of subsets of X that satisfies the following requirements

1. $\emptyset, X \in \mathcal{T}$.
2. $\mathcal{T}$ is closed under union: any (possibly infinite) union $\cup_{j \in J} U_j$ of elements $U_j \in \mathcal{T}$ is again an element of $\mathcal{T}$, i.e. $(\cup_{j \in J} U_j) \in \mathcal{T}$.
3. $\mathcal{T}$ is closed under finite intersection: any finite intersection $\cap_{j \in J} U_j$ of elements $U_j \in \mathcal{T}$ is again an element of $\mathcal{T}$, i.e. $(\cap_{j \in J} U_j) \in \mathcal{T}$.

The elements U_i of the topology $\mathcal{T}$ are called the **open sets** of X . The tuple $(X, \mathcal{T})$ of X and $\mathcal{T}$ is called a **topological space**. In this context, the set X is often referred to as the underlying set of $(X, \mathcal{T})$. Note that we will frequently denote the topological space by X as well.

For any given set X there are two particularly simple topologies. The first one is $\mathcal{T}_0 := \{X, \emptyset\}$ consisting only of the set itself and the empty set. It is called the trivial topology. The second one is the power set $\mathcal{T}_\Omega := \Omega_X$ of X , i.e. the collection of all subsets of X , and called the discrete topology. Another important example arises when we turn to a class of spaces with which we are quite familiar: metric spaces. To see how these form a subset of topological spaces, recall that a metric is a function $d : X \times X \rightarrow \mathbb{R}$ that is (i) symmetric $d(x, y) = d(y, x)$, (ii) positive definite, $d(x, y) \geq 0$ where equality holds iff $x = y$, and (iii) satisfies the triangle inequality $d(x, z) \leq d(x, y) + d(y, z)$. A metric space is then an ordered pair (X, d) of an underlying set X endowed with a metric d . Importantly, every metric space (X, d) is also a

¹At least as far as scientific publications are concerned.

topological space $(X, \mathcal{T}_d)$ with the topology $\mathcal{T}_d$ of open balls

$$B_r(x) = \{y \in X \mid d(x, y) < r\} \quad \forall x \in X \quad (2.1)$$

and all their possible unions. The metric axioms that d satisfies ensure that $\mathcal{T}_d$ is indeed a topology as defined in Def. 2.1.1. A topology $\mathcal{T}_d$ that is defined in this way by a metric d is called a metric topology. Every metric space is therefore also a topological space. The converse is not true: a topological space does not necessarily support a compatible metric. Thus, the class of metric spaces forms a strict subset of the class of topological spaces.

Let us briefly mention that there is a natural way of defining topological subspaces. Given a topological space $(X, \mathcal{T})$ and any subset $S \subseteq X$. We can endow S with the subspace topology

$$\mathcal{T}_S := \{S \cap U \mid U \in \mathcal{T}\} \quad (2.2)$$

to make it into a topological subspace $(S, \mathcal{T}_S)$ of $(X, \mathcal{T})$. Properties of a topological space $(X, \mathcal{T})$ that are inherited by topological subspaces $(S, \mathcal{T}_S)$ are called *hereditary*.

2.1.2 Continuity

At the beginning of this section, we described topology as rubber-sheet geometry. In more mathematical terminology, one would say that topology is geometry up to continuous transformation. In physics it is quite natural to think of continuity in terms of ϵ - δ continuity which tells us that a function $f : X \rightarrow Y$ is continuous at $x_0 \in X$ if

$$\forall \epsilon > 0 \exists \delta > 0 \mid d_X(x, x_0) < \delta \implies d_Y(f(x), f(x_0)) < \epsilon \quad \forall x \in X. \quad (2.3)$$

Note that this is naturally based on some metrics d_X and d_Y of the underlying spaces X and Y . Having established that a topological space need not be metric, we require a purely topological definition of continuity. This definition is remarkably simple [39].

Definition 2.1.2. Let $(X, \mathcal{T}_X)$ and $(Y, \mathcal{T}_Y)$ be topological spaces. A map $f : X \rightarrow Y$ is **continuous** if the inverse image of an open set in Y is an open set in X , i.e. if $\forall A \in \mathcal{T}_Y \exists B \in \mathcal{T}_X \mid B = f^{-1}(A)$.

Importantly, Def. 2.1.2 reproduces ϵ - δ -continuity for maps between topological spaces carrying a metric topology. Let us consider some examples to illustrate how non-metric topologies work with this definition of continuity.

Example 2.1.1. Let $(X, \mathcal{T}_X)$ and $(Y, \mathcal{T}_Y)$ be topological spaces and let $f : X \rightarrow Y$ be a function. There are the following special cases.

1. If f is constant, i.e. if $f(x) = y$ for all $x \in X$ and some $y \in Y$, then f is continuous for all possible topologies $\mathcal{T}_X$ on X and $\mathcal{T}_Y$ on Y because for any open set $U \in \mathcal{T}_Y$ we either have $f^{-1}(U) = \emptyset$ if $y \notin U$ or $f^{-1}(U) = X$ if $y \in U$, both of which are always open in any topology $\mathcal{T}_X$ on X .
2. If $\mathcal{T}_X = \Omega_X$, i.e. if X carries the discrete topology, then f is continuous regardless of the topology $\mathcal{T}_Y$ on Y because for any open set $U \in \mathcal{T}_Y$ the preimage $f^{-1}(U) \in \mathcal{T}_X = \Omega_X$ by definition of the power set.
3. If $\mathcal{T}_Y = \{\emptyset, Y\}$, i.e. if Y carries the trivial topology, then f is continuous regardless of the topology $\mathcal{T}_X$ on X because the only open sets on Y are $\emptyset$ and Y the preimages of which are $f^{-1}(\emptyset) = \emptyset$ and $f^{-1}(Y) = X$, both of which are always open in any topology $\mathcal{T}_X$ on X .
4. If $\mathcal{T}_Y = \Omega_Y$ or if $\mathcal{T}_X = \{\emptyset, X\}$, then the only functions that are always continuous are constant functions (see above).

As a more concrete example, consider the function $f : X \rightarrow Y, x \mapsto x^2$ where $X = Y = \mathbb{R}$. We know that f is continuous when X and Y are both equipped with the standard metric topology $\mathcal{T}_d$ of $\mathbb{R}$. For instance, the inverse image of the open interval $(0, 1) \subset Y$ is the set $(-1, 0) \cup (0, 1) \subset X$, which is open in $\mathcal{T}_d$ on X even though it appears to be split at zero. Importantly, this example demonstrates why the naive definition of continuity – “a map f is continuous if it maps an open set to an open set” – does not work: the open interval $(-1, 1)$ in X is mapped to the half-closed interval $[0, 1)$ in Y , which is not open in the standard metric topology on Y . Note, that according to Ex. 2.1.1, even the familiar function $f(x) = x^2$ would cease to be continuous if we equipped $Y = \mathbb{R}$ with the discrete topology or $X = \mathbb{R}$ with the trivial topology, while keeping the metric topology for the respective other copy of $\mathbb{R}$.

2.1.3 Separation

In order to further analyse the proximity structure that the open sets of a topology impose on an underlying set, we require some additional concepts. Let $(X, \mathcal{T})$ be a topological space and let $S \subseteq X$ be a subset of X . Then S is said to be *closed* if its complement is open, i.e. if $(X - S) \in \mathcal{T}$. Note that according to this definition, X and $\emptyset$ are both open and closed,² so openness and closedness of a set are not mutually exclusive properties. Furthermore, whether a set $S \subseteq X$ is open or closed depends not only on the topology $\mathcal{T}$ but also on the underlying set X . For any (open or closed) subset $S \subseteq X$ we can continue to define its closure $\bar{S}$ as the smallest closed set containing S . Similarly, we define the interior $\mathring{S}$ of S as the largest open subset contained in S . The boundary $b(S)$ of S is then the complement of $\mathring{S}$ in S , i.e. $b(S) := (S - \mathring{S})$. These definitions already give us a rudimentary idea of how the subsets of a topological space relate to one another. To further refine this idea, we follow Ref. [39] and introduce the concept of neighbourhoods first.

Definition 2.1.3. Let $(X, \mathcal{T})$ be a topological space and let $S \subseteq X$ be a subset in X . A **neighbourhood** of S is a subset $N \subseteq X$ that includes an open set $U \in \mathcal{T}$ that contains S , that is

$$S \subseteq U \subseteq N \subseteq X.$$

This is equivalent to saying that S is in the interior N° of N . A neighbourhood need not be an open set itself. If it is, we call it an **open neighbourhood**.

Intuitively speaking, a neighbourhood N of a set S is a collection of points that contains S , but still allows one to move away from S by a certain amount in any direction without leaving N . In this sense, the family of all neighbourhoods of a given subset S indicates how well the topology can distinguish it. In the trivial topology every non-trivial subset has precisely one neighbourhood, namely the entire space, so it is essentially incapable of distinguishing non-trivial subsets. In particular, it cannot distinguish between any two distinct points $p_1 \neq p_2$ of X . In a discrete topological space, on the other hand, every non-trivial subset can be distinguished arbitrarily well. Accordingly, the discrete topology even allows for perfect distinction between any two distinct points. Since subsets of individual points represent the smallest possible class of non-trivial subsets, the ability of a topology to distinguish between them is of particular importance in general topology. In fact, it is the concept of point distinguishability that forms the basis of the separation axioms – a comprehensive framework for assessing the distinguishing capacities of topologies. In order to explore the idea of point distinguishability and separation axioms further, we first introduce a formal version of the topological distinguishability of points.

Definition 2.1.4. Let $(X, \mathcal{T})$ be a topological space and let $p_1, p_2 \in X$ be points in X with neighbourhoods $N_i = \{N \subseteq X \mid N \text{ is neighbourhood of } p_i\}$. Then p_1 and p_2 are said to be **topologically indistinguishable**, $p_1 \equiv p_2$, if and only if they have the same neighbourhoods $N_1 = N_2$.

Accordingly, two points $p_1, p_2 \in X$ of a topological space $(X, \mathcal{T})$ that do not have the same neighbourhoods are called topologically distinguishable. In that case, at least one of the points has a neighbourhood that is not a neighbourhood of the other point. Note that it is entirely possible for both points to possess such exclusive neighbourhoods, but it is not required to be so by topological distinguishability alone. Tightening the conditions in this respect leads to the definition of separability of points.

Definition 2.1.5. Let $(X, \mathcal{T})$ be a topological space and let $p_1, p_2 \in X$ be points in X with neighbourhoods $N_i = \{N \subseteq X \mid N \text{ is neighbourhood of } p_i\}$. Then p_1 and p_2 are said to be **separable** if there exist neighbourhoods $n_1 \in N_1$ and $n_2 \in N_2$ such that

$$n_1 \not\subseteq N_2 \quad \text{and} \quad n_2 \not\subseteq N_1. \quad (2.4)$$

Note that Def. 2.1.5 naturally extends to larger subsets $S_1, S_2 \subseteq X$ of the underlying set X , i.e. one can replace the points p_1, p_2 in Def. 2.1.5 by any subsets S_1, S_2 to define their separability. However, we will focus on point separability here. Based on Def. 2.1.5 one can conceive increasingly powerful types of point separability. For example, two points $p_1, p_2 \in X$ are even separable by neighbourhoods if they have disjoint neighbourhoods, i.e. neighbourhoods N_1 of p_1 and N_2 of p_2 such that $N_1 \cap N_2 = \emptyset$. Now

²Such subsets are sometimes called *clopen*.

	$p_1 \not\equiv p_2$	$p_1 \neq p_2$
S	symmetric	Fréchet
SN	preregular	Hausdorff

Table 2.1: Names of topological spaces allowing the separation (S) or separation by neighbourhoods (SN) of any two distinct ($p_1 \neq p_2$) or topologically distinguishable ($p_1 \not\equiv p_2$) points.

one can ask the following kind of questions. Are any two distinct points $p_1 \neq p_2$ in X separable? How about separable by neighbourhoods? How about any two topologically distinguishable points $p_1 \not\equiv p_2$, are those separable or even separable by neighbourhoods? Each combination of such properties is called a separation axiom and has, of course, its own name. The ones outlined before are summarised in Tab. 2.1. Further separation axioms arise when one applies the separability conditions not to pairs of distinct and topologically distinguishable points but to disjoint pairs made up of one point and one closed set or disjoint pairs of closed sets.

Of the many conceivable separation axioms, the Hausdorff axiom is by far the most frequently used. This is because a topological space satisfying the Hausdorff axiom is, in many ways, a *nice* space: its topology can distinguish points in a strong sense, ensuring they are not too close together. This imposes a certain regularity on the space. For instance, it implies the uniqueness of important objects like limits and sequences. For this reason, the Hausdorff axiom is often required in topological construction schemes and proofs. A topological space that satisfies the Hausdorff axiom is called a Hausdorff space and the Hausdorff property is hereditary. Importantly, all metric spaces are Hausdorff spaces, such that the Hausdorff criterion is readily satisfied in most physics contexts.

2.1.4 Connectedness

Earlier, in the discussion of $f(x) = x^2$ and its continuity, we found that $f^{-1}((0, 1)) = (-1, 0) \cup (0, 1)$. We argued that this set is open in the standard metric topology on $\mathbb{R}$ even though it is clearly different from a naive open set like $(-1, 1)$ in some way. Visually, the difference between $(-1, 1)$ and $(-1, 0) \cup (0, 1)$ is that the former is connected, whereas the latter is not. In fact, they differ with respect to a fundamental topological concept that is, in line with this simple observation, called connectedness [39].

Definition 2.1.6. Let $(X, \mathcal{T})$ be a topological space. Then X is said to be **connected** if it cannot be written as $X = X_1 \cup X_2$ with non-trivial open sets X_1, X_2 fulfilling $X_1 \cap X_2 = \emptyset$. Otherwise it is called **disconnected**.

The notion connectedness depends heavily on the choice of topology. For instance, in the trivial topology, every space is connected because the only non-trivial open set is all of X , whereas in the discrete topology, every space is totally disconnected because every single point of X is open by itself. There is a slightly stronger notion of connectedness that is based on continuous functions.

Definition 2.1.7. Let $(X, \mathcal{T})$ be a topological space. Then X is said to be **path connected** if for any points $p_0, p_1 \in X$ there exists a continuous function $f : [0, 1] \rightarrow X$ such that $f(0) = p_0$ and $f(1) = p_1$.

Path-connectedness always implies connectedness, while the converse is not true, making path-connectedness a stronger notion than connectedness. Accordingly, the topological subspace $(-1, 0) \cup (0, 1)$ from before is indeed disconnected by definition. This shows that connectedness is not a hereditary property.

2.1.5 Compactness

So far we have discussed how topology organises a space in terms of proximity, separability and connectedness. Another important aspect of a topological space is its *extent*, which is captured by a fundamental notion called compactness. The meaning of compactness can be difficult to appreciate, which is why we will explore its basic idea before presenting its formal definition. Again, we follow Ref. [39] and start with the definition of a cover.

Definition 2.1.8. Let $(X, \mathcal{T})$ be a topological space. A family A of subsets of X is called a **cover** of X if it fulfils

$$X = \bigcup_{S \in A} S. \quad (2.5)$$

If all sets of a cover A are open sets then A is said to be an open cover. We are predominantly interested in such open covers. It is not particularly difficult to come up with a simple example: the family $A_0 = \{X\}$ constitutes a trivial open cover of any topological space X . However, A_0 is so simple that we can hardly learn anything worth knowing from it. For example, it is impossible for us to decide whether X is connected based on A_0 alone. This is because X is disconnected if and only if there exists an open cover that splits into disjoint subcovers and the trivial cover A_0 cannot be split in this way. In order to test connectedness we would therefore need more complex covers consisting of larger numbers of subsets. In particular, we would have to examine all possible covers to make a final decision. Intuitively, the analysis of connectedness in a topological space with n connected components would only really require covers containing n subsets and all the finer covers are redundant. However, depending on the topological space under consideration, the investigation of some properties may even require covers consisting of an infinite number of subsets. The idea of compactness is to associate the extent of a topological space with the need for such infinitely large covers: a space is called compact if every question about its properties can be answered in terms of finite covers. The formal definition reads as follows.

Definition 2.1.9. Let $(X, \mathcal{T})$ be a topological space and let $\mathcal{C}$ be the collection of all covers of X . Then X is said to be **compact** if every open cover of X has a finite subcover, i.e. if for every cover A there exists a **finite** subcover F such that

$$X = \bigcup_{C \in F} C. \quad (2.6)$$

In many ways, compact topological spaces can be thought of as *finite-like* spaces and they have many desirable properties reflecting this finiteness. For instance, continuous functions on compact spaces always attain their extrema, and every infinite sequence of points in a compact space has a limit point within the space. The latter is the reason why compact spaces are sometimes described as spaces where there is “nowhere to escape to”. A subset $Y \subset X$ of a topological space is compact if it is compact as a subspace in the subspace topology. That is, Y is compact if for every family A of open subsets of X that covers Y , i.e. $Y \subseteq \bigcup_{S \in A} S$, there exists a finite subfamily $F \subseteq A$ that still covers Y , i.e. $Y \subseteq \bigcup_{S \in F} S$. With this we can see that individual points $p \in X$ are always compact since every cover can be reduced to a cover consisting of just one single open subset U with $p \in U \subseteq X$. In Hausdorff spaces, compact subsets are necessarily closed. One of the most important examples of this is $\mathbb{R}^n$ where subsets are compact if and only if they are closed and bounded, as stated in the famous Heine–Borel theorem. There is another interesting interplay between compactness and Hausdorff-ness: if a topological space $(X, \mathcal{T})$ is compact and Hausdorff, then no finer topology on X is compact and no coarser topology on X is Hausdorff.

2.1.6 Homeomorphisms

The idea of topology as geometry up to continuous deformation raises another important question: how can we determine whether two topological objects can be continuously deformed into one another? In more technical terms, this evolves into the much bigger question of how we can partition the set of all topological spaces into equivalence classes of spaces that are topologically identical. To address this problem, we need one final concept, namely that of a homeomorphism [39].

Definition 2.1.10. Let X and Y be topological spaces. A map $f : X \rightarrow Y$ is called a **homeomorphism** if it is continuous and has an inverse $f^{-1} : Y \rightarrow X$ that is continuous.

If there is a homeomorphism between two topological spaces X and Y , they are said to be homeomorphic. The notion of homeomorphisms enables the definition of many fundamental objects. One of the most relevant for theoretical physics is that of a topological manifold.

Definition 2.1.11. An **n-dimensional topological manifold** M is a topological Hausdorff space $(X, \mathcal{T})$ that is locally Euclidean. A topological space is locally Euclidean if each point in X has an open neighbourhood that is homeomorphic to an open subset of $\mathbb{R}^n$, i.e. if for each $x \in X$ we can find

1. an open neighbourhood $U \in \mathcal{T}$ with $x \in U$,
2. an open subset $\hat{U} \subseteq \mathbb{R}^n$, and
3. a homeomorphism $\varphi : U \rightarrow \hat{U}$.

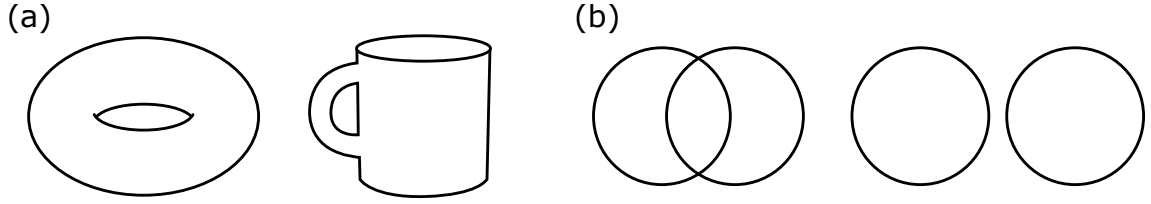


Figure 2.1: Examples of homeomorphic spaces: (a) doughnut and cup, (b) interlinked and separated rings. Illustration created by the author, inspired by Ref. [39].

Note that it is not unusual to include additional requirements, such as compactness, in the definition of a topological manifold. The definition we give above can be considered the minimal definition of a topological manifold that most authors agree on. Manifolds generalise Euclidean geometry to less regular objects. Importantly, these need not be embedded in a higher-dimensional Euclidean space. The fact that a manifold is locally Euclidean means that we can borrow many of the established notions about calculus from Euclidean space. If this is done in a consistent fashion, one arrives at the definition of a C^k differentiable structure which implements differential calculus of k -times differentiable functions on a given manifold. A differentiable manifold is then a topological manifold with a compatible differentiable structure. The manifolds we encounter in theoretical physics are usually smooth differentiable manifolds with a C^∞ differentiable structure.

Being homeomorphic provides an equivalence relation $\sim$ on the set $\mathfrak{X} = \{X \mid X \text{ is a topological space}\}$ of all topological spaces, as it naturally satisfies reflexivity $X \sim X$, symmetry $X \sim Y \Rightarrow Y \sim X$, and transitivity $X \sim Y, Y \sim Z \Rightarrow X \sim Z$.³ Picture (a) of Fig. 2.1 shows the most prominent example of topological equivalence: a coffee cup is equivalent to a doughnut. Picture (b) shows a less intuitive example: the separated rings are homeomorphic to the interlinked rings. Our intuition tends to suggest that these cannot be deformed into one another without cutting and glueing. However, this is due to our habit of visualising objects as embedded in $\mathbb{R}^3$, where such a homeomorphism does indeed not exist. However, embeddings of the separated and interlinked rings in $\mathbb{R}^3$ are not the same as the separated and interlinked rings themselves: they are only instantiations, or realisations, of the separated and interlinked rings and as such subject to restrictions imposed by the embedding itself.⁴ This failure of intuition is a recurring theme in topology and serves as a cautionary tale for us. It also illustrates that it is impractical to go over all topological spaces and try to put down pairwise homeomorphisms let alone a proof of their non-existence for topologically distinct spaces. How then can we characterise the homeomorphism classes of topological spaces? The fact that a complete answer to this question is not yet known tells us how hard a problem it is.

Facing the overwhelming difficulty of the complete homeomorphism classification, we turn to a more feasible strategy. By shifting our focus to quantities of topological spaces that remain unchanged under homeomorphisms, we can at least identify situations in which two topological spaces *cannot* be homeomorphic. Quantities that remain invariant under homeomorphisms are called topological invariants and even though a distinction based on topological invariants is a far less profound classification scheme than the full homeomorphism classification, it is extremely useful in practice.

2.1.7 (The Hole Truth About) Topological Invariants

One can show that the topological properties that we have discussed so far – connectedness, compactness or being Hausdorff – are invariant under homeomorphisms. Accordingly, they serve as simple examples of topological invariants. Let us provide a generic definition.

Definition 2.1.12. A **topological invariant** of a topological space $(X, \mathcal{T})$ is any quantity of the underlying space that is left invariant by homeomorphisms.

³Reflexivity is ensured by the trivial homeomorphism $f = \text{id}_X$, symmetry is fulfilled automatically because if $f : X \rightarrow Y$ is a homeomorphism then so is $f^{-1} : Y \rightarrow X$ by definition, and transitivity follows from the fact that the composition $g \circ f : X \rightarrow Z$ of homeomorphisms $f : X \rightarrow Y$ and $g : Y \rightarrow Z$ inherits the homeomorphism properties from f and g .

⁴There are other familiar examples of such embedding artefacts: the self-penetration of the Klein bottle in $\mathbb{R}^3$ or the non-vanishing local curvature of the two-torus in $\mathbb{R}^3$.

If we knew *all* the topological invariants of a given topological space, we could specify its homeomorphism equivalence class. However, as long as we only now a partial list of topological invariants, the best we can do is the following:

Two topological spaces that differ in at least one topological invariant cannot be homeomorphic.

Strikingly, there exist topological invariants that take the form of a mere number. For a disconnected topological space this might be the number of connected components. However, there are much more profound *number invariants* than this. For instance, being locally Euclidean is a topological invariant which makes the dimension d of a manifold a topological number invariant too.

Another famous topological invariant of this type was discovered by Riemann when he was thinking about connectedness: he wondered how many types of one-dimensional simple,⁵ non-intersecting closed curves $\gamma^1 : \mathbb{S}^1 \rightarrow M$ could be simultaneously removed from a two-dimensional manifold M without disconnecting it [62]. The result leads to the *number of holes* in M that we call the genus of M today.

The notion of detecting and utilising the holes of a topological space to classify it became a popular approach. Betti, for example, realised that one-dimensional simple closed curves can only detect holes of codimension one in M . In particular, they cannot capture zero-dimensional holes in three dimensions. Based on this realisation he expanded the concept to include the n -dimensional equivalent of simple closed curves, i.e. injective and continuous maps $\gamma^n : \mathbb{S}^n \rightarrow M$. The number of distinct equivalence classes of n -dimensional simple closed curves that can be removed from a space without disconnecting it is called the n -th Betti number β_n . Note that the simple closed curves that give rise to the Betti numbers may intersect, whereas the curves that generate the genus may not.

The basic idea of holes in topological spaces inspired the development of an entirely new field of mathematics called algebraic topology. Roughly speaking, algebraic topology uses algebraic tools to study topological spaces. This allows one to recast statements about topological spaces into statements about algebraic structures like groups. These come with a lot of extra structure and often make claims easier to prove and results easier to grasp. Among the major disciplines of algebraic topology are homotopy and (co-)homology theory. For certain types of topological spaces, algebraic topology provides us with a particularly beautiful framework for generating topological number invariants. These topological spaces are known as fibre bundle spaces and their topological number invariants are called characteristic numbers. We will cover some of the fundamental ideas of algebraic topology in the following and we will come back to the notion of fibre bundles and their characteristic numbers in a later section. In particular, we will encounter a class of complex bundle spaces that naturally arise in the mathematical description of certain condensed matter systems. These spaces are associated with a characteristic number that became famous for its immediate physical significance: the Chern number of a ground state bundle, which represents the quantised Hall conductance [10].

2.1.8 Homotopy

We are closely following Ref. [39]. The genus of a two-dimensional manifold M is the maximum number of one-dimensional simple, non-intersecting closed curves that can be simultaneously removed from M without disconnecting it. In some simple cases, this imagery is enough for us to obtain a concrete result. For example, Fig. 2.2 shows three simple closed curves a, b and c on the two-sphere $\mathbb{S}^2$ and the two-torus $\mathbb{T}^2$, respectively. We see that removing any one of the three loops from $\mathbb{S}^2$ will disconnect it, while $\mathbb{T}^2$ remains connected if a or b is removed. Interestingly, the removal of a (b) from $\mathbb{T}^2$ cuts b (a), such that it is no longer a closed curve afterwards. In fact, the only closed curves that remain are of type c and will disconnect the space upon their removal. Accordingly, we count a maximum number of $g_{\mathbb{S}^2} = 0$ loops that can be removed without disconnecting $\mathbb{S}^2$ and a maximum number of $g_{\mathbb{T}^2} = 1$ loop that can be removed without disconnecting $\mathbb{T}^2$. Clearly, these values for the genera of $\mathbb{S}^2$ and $\mathbb{T}^2$ coincide with the numbers of holes we count if we just stare at the geometrical objects. However, there seem to be two distinct types of loops, namely a and b type loops, that can be removed from $\mathbb{T}^2$ without disconnecting it. The fact that this observation does not appear to be accounted for by the genus $g_{\mathbb{T}^2} = 1$ indicates that there is more to learn here. The natural way to approach this is to ask whether a and b are actually

⁵A continuous curve is called simple if it is not self-intersecting.

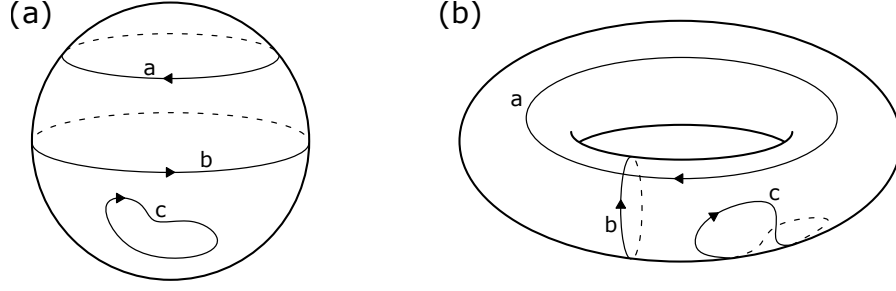


Figure 2.2: Different closed curves in (a) a two-sphere $\mathbb{S}^2$ and (b) a two-torus $\mathbb{T}^2$. Illustrations created by the author, based on Refs. [63, 64].

distinct curves in $\mathbb{T}^2$. If we reinstate the rubber-sheet picture we used before and apply it to the loops, we find that no amount of pushing and pulling is going to deform loop a into loop b within $\mathbb{T}^2$ – they *do* seem to be distinct. In fact, neither loop a nor loop b can be deformed into loop c within $\mathbb{T}^2$ either. Consequently, all three types of loops seem to be distinct. In contrast, any pair of loops in $\mathbb{S}^2$ can be deformed into one another. What we need in order to investigate this idea further is a formal way of comparing pairs of loops $f, g : \mathbb{S}^1 \rightarrow X$. More generally, we are interested in a comparison of continuous maps between any two topological spaces X and Y . The missing concept is that of a homotopy [39].

Definition 2.1.13. Let $(X, \mathcal{T}_X)$ and $(Y, \mathcal{T}_Y)$ be topological spaces and let $g, h : X \rightarrow Y$ be continuous functions. A **homotopy** between g and h is a continuous map

$$F : X \times [0, 1] \rightarrow Y$$

$$(x, t) \mapsto F(x, t) = f_t(x), \quad (2.7)$$

such that $F(x, 0) = f_0(x) = g(x)$ and $F(x, 1) = f_1(x) = h(x)$. The family $\{f_t\}_{t \in [0, 1]}$ of continuous maps $f_t : X \rightarrow Y$ is called a continuous family of maps.

The existence of a homotopy tells us that $f_0 : X \rightarrow Y$ and $f_1 : X \rightarrow Y$ can be continuously deformed into one another. Two continuous maps $g, h : X \rightarrow Y$ between topological spaces X and Y are said to be homotopic, $g \sim h$, if there exists a homotopy between them. Being homotopic is an equivalence relation on the continuous functions from X to Y . The homotopy equivalence class of a continuous function $f : X \rightarrow Y$ is denoted by $[f] = \{g : X \rightarrow Y \mid f \sim g\}$ and is called its homotopy class. Homotopy allows us to relax the notion of homeomorphism classes of topological spaces and get a coarser but more manageable classification.

Definition 2.1.14. Let $X \equiv (X, \mathcal{T}_X)$ and $Y \equiv (Y, \mathcal{T}_Y)$ be topological spaces. We say X and Y are **homotopy equivalent** or **homotopic**, written as $X \sim Y$, if there exist continuous maps $f : X \rightarrow Y$ and $g : Y \rightarrow X$ such that $(f \circ g) \sim \text{id}_Y$ and $(g \circ f) \sim \text{id}_X$.

Being homeomorphic is a stronger notion than being homotopic. For example, a point p and the real line $\mathbb{R}$ are of the same homotopy type, whereas p is clearly not homeomorphic to $\mathbb{R}$. A topological space X that is homotopic to a point is said to be *contractible*. Being contractible is an example of a topological homotopy invariant. Note that homotopy invariants form a *stronger* class of topological invariants than homeomorphism invariants, precisely *because* the homotopy type provides a *weaker* equivalence relation than the homeomorphism type. Every topological homotopy invariant is therefore a topological homeomorphism invariant, while the converse is not necessarily the case. In order to use homotopy theory to explore the notion of topological holes, we recall the formal definition of loops.

Definition 2.1.15. Let $I = [0, 1]$ denote the unit interval and let $X = (X, \mathcal{T}_X)$ be a topological space. Any map $\alpha : I \rightarrow X$ with $\alpha(0) = \alpha(1) = x_0 \in X$ is called a **loop** at x_0 . Due to the identification $\alpha(0) = \alpha(1)$ this is equivalent to a map $\alpha : \mathbb{S}^1 \rightarrow X$.

Of course, loops are just special classes of maps between topological spaces so we can partition them according to their homotopy classes. Given any loop α in X , its homotopy class is denoted by $[\alpha]$. We will see that the set of homotopy classes of loops acquires a group structure when it is endowed with the product of path concatenation.

Definition 2.1.16. Let $(X, \mathcal{T}_X)$ be a topological space. The set of homotopy classes of loops

$$\{[\alpha_j]_{j \in J} \mid \alpha_j : I \rightarrow X \text{ with } \alpha_j(0) = \alpha_j(1) = x_0\} \quad (2.8)$$

at any given point $x_0 \in X$ together with the path concatenation has group structure and is called the **fundamental group** $\pi_1(X, x_0)$ of X based at x_0 .

In order to convince ourselves that the fundamental group has a group structure under path concatenation we first define the concatenation of homotopy classes as

$$[\alpha] * [\beta] = [\alpha * \beta]. \quad (2.9)$$

We note that Eq. (2.9) is well-defined; it does not depend on the choice of representatives. The set of loops at x_0 is also closed under loop concatenation because any concatenation of loops $\alpha, \beta : I \rightarrow X$ with $\alpha(0) = \alpha(1) = \beta(0) = \beta(1) = x_0$ yields

$$(\alpha * \beta) : I \rightarrow X \text{ with } x \mapsto \begin{cases} \alpha(2x) & x \leq \frac{1}{2} \\ \beta(2x) & x > \frac{1}{2}, \end{cases} \quad (2.10)$$

which clearly describes another loop at x_0 . The associativity of $*$ follows from the associativity of map concatenation and the identity element is given by the constant map $c_{x_0}(x) = x_0$. Finally, the inverse of a loop $\alpha(t)$ is simply $\alpha^{-1}(t) \equiv \alpha(1 - t)$, i.e. the same loop traversed in the opposite direction.

In general, the fundamental group is a local property: it is attached to every base point $x_0 \in X$ independently. However, there is a way to get rid of this locality in some cases. To see this, consider two base points $x_0, x_1 \in X$ with fundamental groups $\pi_1(X, x_0)$ and $\pi_1(X, x_1)$ and suppose that there exists a continuous path $h : I \rightarrow X$ with $h(0) = x_0$ and $h(1) = x_1$. Then h induces a map

$$\begin{aligned} \eta_h : \pi_1(X, x_1) &\rightarrow \pi_1(X, x_0) \\ [f] &\mapsto [(h * f) * h^{-1}], \end{aligned} \quad (2.11)$$

which is a group isomorphism with inverse map $\eta_{h^{-1}} \equiv \eta_h^{-1} : \pi_1(X, x_1) \rightarrow \pi_1(X, x_0)$. Thus, in any path-connected topological space the fundamental group is independent of the base point and we can simply write

$$\pi_1(X, x_0) \equiv \pi_1(X) \quad (2.12)$$

for all $x_0 \in X$. In a disconnected space the fundamental groups of its connected components may very well differ. Finally, we call a topological space X *simply connected* if it is path-connected and has a trivial fundamental group $\pi_1(X) = \{e\}$. Importantly, the fundamental group is invariant under homotopies which makes it our first example of an *algebraic* topological homotopy invariant.

Let us consider an important example: the fundamental group $\pi_1(\mathbb{S}^1)$ of the circle. The basic idea for constructing it is that every loop $\alpha : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ can be uniquely *lifted* to $\mathbb{R}$. To show this, we take

$$\mathbb{S}^1 = \{z \in \mathbb{C} \mid |z| = 1\} \subset \mathbb{C}, \quad (2.13)$$

and define a projection map

$$\begin{aligned} p : \mathbb{R} &\rightarrow \mathbb{S}^1 \subset \mathbb{C} \\ x &\mapsto e^{ix}, \end{aligned} \quad (2.14)$$

that wraps $\mathbb{R}$ around $\mathbb{S}^1 \subset \mathbb{C}$ as illustrated in Fig. 2.3. Note that any $x, y \in \mathbb{R}$ fulfilling $x - y = 2\pi n$ for some $n \in \mathbb{Z}$ are mapped to the same point in $\mathbb{S}^1$ by the projection map. Such pairs of points are therefore equivalent with respect to p and we can write $x \sim y$. Consequently, the equivalence class

$$[x] = \{y \in \mathbb{R} \in \mathbb{C} \mid x - y = 2\pi n, n \in \mathbb{Z}\} \quad (2.15)$$

is identified with the same point $p(x) = e^{ix}$ in $\mathbb{S}^1$. Equation (2.15) tells us that a curve

$$\begin{aligned} \omega : I &\rightarrow \mathbb{S}^1 \\ s &\mapsto e^{i2\pi f(s)} \end{aligned} \quad (2.16)$$

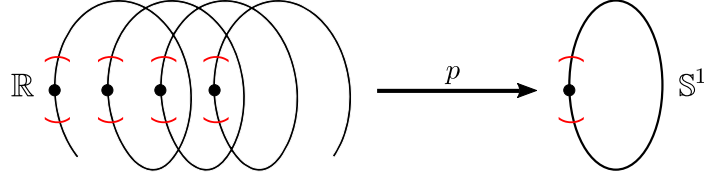


Figure 2.3: Sketch of $\mathbb{R}$ as a covering space of $\mathbb{S}^1 \subset \mathbb{C}$ by the projection map p . Black dots and red intervals in $\mathbb{R}$ mark the preimages of the black dot and red interval in $\mathbb{S}^1$ under p .

is closed if and only if $f(s)$ satisfies $f(0) = 0$ and $f(1) = n$ for some $n \in \mathbb{Z}$. Accordingly, every loop can be uniquely labeled by an integer $n \in \mathbb{Z}$ and its equivalence class $[\omega_n]$ can be represented by the simple map $\omega_n(s) = e^{i2\pi ns}$ where $f(s) = ns$. Since $\mathbb{S}^1 \subset \mathbb{C}$ is path-connected, the map

$$\begin{aligned} \Phi : \mathbb{Z} &\rightarrow \pi_1(\mathbb{S}^1, 1) \\ n &\mapsto [\omega_n] \end{aligned} \quad (2.17)$$

from the integers to the fundamental group $\pi_1(\mathbb{S}^1, 1)$ at $x_0 = 1$ defines a group isomorphism between $\mathbb{Z}$ and $\pi_1(\mathbb{S}^1)$, giving

$$\pi_1(\mathbb{S}^1) \simeq \mathbb{Z}. \quad (2.18)$$

When we analysed the genera of $\mathbb{S}^2$ and $\mathbb{T}^2$ before, we speculated that there might be more to learn from the different types of closed curves because the genus of $\mathbb{T}^2$ did not account for the two distinct ways in which the torus tolerated the removal of either loop a or loop b in Fig. 2.2. Does the fundamental group resolve this missing information? We established earlier that there is only a single equivalence class of loops in $\mathbb{S}^2$ so its fundamental group is bound to be trivial, i.e. $\pi_1(\mathbb{S}^2) \simeq \{e\}$. In order to determine the fundamental group of the two-torus we use the following theorem about topological product spaces.

Theorem 2.1.1. *Let $X = (X, \mathcal{T}_X)$ and $Y = (Y, \mathcal{T}_Y)$ be path-connected topological spaces. Then the fundamental group of the product space $X \times Y$ is isomorphic to the direct sum of the fundamental groups of the individual topological spaces, i.e.*

$$\pi_1(X \times Y) \simeq \pi_1(X) \oplus \pi_1(Y).$$

We can use this to determine $\pi_1(\mathbb{T}^2)$ by writing $\mathbb{T}^2 \simeq \mathbb{S}^1 \times \mathbb{S}^1$ which immediately gives

$$\pi_1(\mathbb{T}^2) \simeq \pi_1(\mathbb{S}^1 \times \mathbb{S}^1) \simeq \pi_1(\mathbb{S}^1) \oplus \pi_1(\mathbb{S}^1) \simeq \mathbb{Z} \oplus \mathbb{Z} = \mathbb{Z}^2. \quad (2.19)$$

Importantly, $\pi_1(\mathbb{T}^2)$ features two copies of $\mathbb{Z}$, which *does* seem to reflect the existence of two distinct types of closed curves that can be removed without disconnecting it. What is more, we find that the first Betti number $\beta_1(\mathbb{T}^2)$, i.e. the maximum number of simple and possibly intersecting closed curves that can be simultaneously removed from $\mathbb{T}^2$ without disconnecting it, is precisely two, namely the type a and type b curves in Fig. 2.2. This is no coincidence, as the first Betti number $\beta_1(X)$ of a two-dimensional closed *orientable* topological manifold M can be defined as the rank⁶

$$\beta_1(M) = \text{rank}(\pi_1(M)), \quad (2.20)$$

of the fundamental group $\pi_1(M)$, cf. e.g. Ref. [65]. Consequently, there exists the general relation

$$\text{rank}(\pi_1(M)) = \beta_1(M) = 2g(M) \quad (2.21)$$

between the fundamental group $\pi_1(M)$, the first Betti number $\beta_1(M)$ and the genus $g(M)$ of any two-dimensional closed orientable manifold M . Indeed, Eq. (2.21) gives the correct genera

$$g(\mathbb{S}^2) = \frac{1}{2} \text{rank}(\{e\}) = \frac{1}{2} \cdot 0 = 0 \quad \text{and} \quad g(\mathbb{T}^2) = \frac{1}{2} \text{rank}(\mathbb{Z} \oplus \mathbb{Z}) = \frac{1}{2} \cdot 2 = 1 \quad (2.22)$$

⁶The rank of a group G is the smallest cardinality of a generating set for G .

where we plugged in the fundamental groups $\pi_1(\mathbb{S}^2) \simeq \{e\}$ and $\pi_1(\mathbb{T}^2) \simeq \mathbb{Z} \oplus \mathbb{Z}$ we have determined before. Equation (2.21) tells us that the fundamental group, the Betti number and the genus ultimately detect the same kind of topological hole in a two-dimensional closed orientable manifold, but with a decreasing level of detail. The fundamental group provides the most refined information, encoding all possible non-trivial homotopy types of loops in terms of directed concatenations of fundamental loops. In contrast, the Betti number only remembers how many different fundamental loop types there are, forgetting about their concatenation. Finally, the genus surrenders even that information in favour of a more visual interpretation: by discarding half the fundamental loop types, it captures the literal number of *handles* that the respective manifold has when it is embedded⁷ in $\mathbb{R}^3$.

What is Wrong with the Fundamental Group?

The fundamental group is great. It is an accessible algebraic topological invariant with a natural graphic interpretation. Still, there are problems that cannot be solved using the fundamental group alone. Take, for example, the problem of distinguishing $\mathbb{S}^2$ from $\mathbb{S}^n$ for $n > 2$. At the moment, we can distinguish $\mathbb{S}^1$ but all of the higher $\mathbb{S}^n$ are simply connected so π_1 does not notice any difference. In order to fix this, it seems obvious to turn to *higher homotopy groups* π_n , i.e. the homotopy classes of maps

$$\omega : \mathbb{S}^n \rightarrow X \quad (2.23)$$

together with the concatenation of maps. Note that the properties of π_1 , which ensure its group structure with concatenation, can easily be generalised to these higher homotopy classes, so that they also have a group structure with concatenation. The higher homotopy groups are therefore well-defined. Handling these groups is not difficult either because it can be shown that *all higher homotopy groups are commutative and thus Abelian*. All of this sounds wonderful – so what is wrong with these higher homotopy groups? There are several answers to this. The first one is that although higher homotopy groups are really easy to define, they are generally very hard to compute. The second, and even more troubling one, is that higher homotopy groups tend to deviate from our geometric expectations. Take, for example, the two-sphere $\mathbb{S}^2$. We have seen that the fundamental group is trivial, $\pi_1(\mathbb{S}^2) \simeq \{e\}$. The second homotopy group $\pi_2(\mathbb{S}^2)$ is, quite reassuringly, again the integers, $\pi_2(\mathbb{S}^2) = \mathbb{Z}$. This seems plausible, given that we found $\pi_1(\mathbb{S}^1) = \mathbb{Z}$ before: the n -th homotopy group seems to detect n -dimensional holes, i.e. holes that can be captured with an n -dimensional sphere. Things look promising, but then $\pi_3(\mathbb{S}^2) = \mathbb{Z}$ is also the integers. How can there be three-dimensional holes in the two-dimensional sphere? This is rather disturbing and it gets even worse. If we go on, we find $\pi_4(\mathbb{S}^2) = \pi_5(\mathbb{S}^2) = \mathbb{Z}_2$, $\pi_6(\mathbb{S}^2) = \mathbb{Z}_{12}$ and, my personal favourite, $\pi_{14}(\mathbb{S}^2) = \mathbb{Z}_{84} \times \mathbb{Z}_2^2$ [66]. So higher homotopy groups are not only hard to compute, but also hard to understand: they simply do not detect what we had hoped they would. Of course, this is a severe problem. One important strategy to circumvent it is provided by the theory of homology.

2.1.9 Homology

One of the biggest problems with higher homotopy groups is that they detect strange high-dimensional holes that cannot be reconciled with our visual conception of holes. In a way, homology groups are designed to fix this: for a d -dimensional topological space X the higher homology groups $H_n(X)$ naturally stop counting at $n = d$. This section is based on Refs. [39] and [61].

The fundamental idea of homology is already present in the guiding principle of *expendable closed curves* that inspired Riemann's and Betti's work on the genus and the Betti numbers. Specifically, it is the realisation that the closed n -dimensional curves, or *n-cycles*, that can safely be removed from a topological space X without disconnecting it, seem to be those that are not themselves boundaries

⁷Note that this is only possible because we assumed orientability. Non-orientable two-dimensional closed manifolds cannot be embedded in $\mathbb{R}^3$ and the whole argument changes. Take, for instance, the Klein bottle $\mathbb{K}^2$. It has a non-trivial fundamental group, $\pi_1(\mathbb{K}^2) \simeq \mathbb{Z} \rtimes \mathbb{Z}$, where the semi-direct product encodes the non-orientability of $\mathbb{K}^2$ and renders $\pi_1(\mathbb{K}^2)$ non-Abelian. Moreover, the *non-orientable* genus $g(\mathbb{K}^2) = 2$ of the Klein bottle is *directly* equal to the rank $\pi_1(\mathbb{K}^2) = 2$ of its fundamental group and does not correspond to one half of the first Betti number, $\beta_1(\mathbb{K}^2) = 1$. The latter is not related to $\text{rank}(\pi_1(\mathbb{K}^2))$ by Eq. (2.20) because $\mathbb{K}^2$ is *not* orientable. This substantially alters the threefold relation from Eq. (2.21). We also observe that the number of handles in the immersion of $\mathbb{K}^2$ into $\mathbb{R}^3$ is most naturally identified with $N = 1$, cf. Fig. 2.7, which is still equal to $\beta(\mathbb{K}^2)$, but no longer directly captured by $g(\mathbb{K}^2)$.

of $(n + 1)$ -dimensional regions in X . Extending this notion, the n -th homology group $H_n(X)$ of a topological space X takes all the n -dimensional closed things in X and then removes those that are boundaries of $(n + 1)$ -dimensional things. This starting point is a fair bit more complicated than the one we assumed for homotopy. By allowing arbitrary n -dimensional cycles, we include objects that are much more general than simple images of n -spheres, say, things like the disjoint union of two simple cycles. However, we also consider a much coarser equivalence relation than being homotopic, namely that of being homologous: specifically, two n -cycles A and B are called *homologous* if they form the boundary ∂W of some $(n + 1)$ -dimensional region W , i.e. if they differ by ∂W . Importantly, two homotopic cycles are always homologous, while two homologous cycles may not be homotopic. A rigorous definition of homology groups will have to formalise the notion of

$$H_n(X) = (n\text{-dim things without boundary}) / (\text{boundaries of } (n + 1)\text{-dim things}), \quad (2.24)$$

where the quotient symbolises a suitable identification of n -cycles with respect to the equivalence relation of being homologous.

There are two main approaches to a homology: simplicial homology and singular homology. Simplicial homology works by partitioning a space X into simple building blocks – simplices like points, line segments, and triangles – which are then organised into a so-called *simplicial complex*. As we will see, this triangulation procedure requires a highly regularised collection of continuous maps from the set of standard simplices into our topological space. Finding such a triangulation can be a challenge, but the extra effort pays off because it makes simplicial homology an intuitive and easy-to-navigate theory that is especially useful for explicit calculations. The problem is, of course, that it is limited to spaces that can be triangulated or reasonably well approximated by simplicial complexes. Since this is not possible for all topological spaces, we need something more. Singular homology generalises the concept of simplicial homology to arbitrary topological spaces by considering unregularised, or *singular*, continuous maps from the set of standard simplices into the respective topological space. Its unregularised nature makes singular homology a versatile and powerful tool for proving things.

The basic constructions of simplicial and singular homology are very similar, so we discuss them for both theories simultaneously in the following. Afterwards, we use simplicial homology to explicitly compute the homology of two simple triangularisable spaces.

Preliminary Considerations

Let us start with some preliminary considerations based on the two-sphere. According to Eq. (2.24), the zeroth homology group $H_0(\mathbb{S}^2)$ of $\mathbb{S}^2$ consists of the zero-dimensional cycles in $\mathbb{S}^2$ modulo those that form boundaries of one-dimensional regions in $\mathbb{S}^2$. In order to make sense of statements like this, we have to impose some algebraic structure first. Namely, we will base the n -th homology group not on the n -dimensional things themselves, but on the free Abelian group generated by them. The free Abelian group $\text{FAb}(\mathcal{B})$ of a collection $\mathcal{B} = \{B_i\}_{i \in I}$ of objects is the Abelian group that has $\mathcal{B}$ for a basis, i.e. it is the group of all formal finite sums of those objects with integer coefficients. Each element $G \in \text{FAb}(\mathcal{B})$ is of the form

$$G = \sum_{n=1}^N g_n B_{i_n}, \quad (2.25)$$

where $N < \infty$, $g_n \in \mathbb{Z}$ and $B_{i_n} \in \mathcal{B}$.⁸ Accordingly, the free Abelian group of any one object $\mathcal{B}_1 = \{B\}$ is isomorphic to the integers, i.e. $\text{FAb}(\mathcal{B}_1) \simeq \mathbb{Z}$. In fact, the free Abelian group of any finite number of N objects $\mathcal{B}_N = \{B_1, \dots, B_N\}$ is isomorphic to the N -fold direct sum of the integers as $\text{FAb}(\mathcal{B}_N) \simeq \mathbb{Z}^N$. Following this, the zeroth homology group of $\mathbb{S}^2$ is based on the free Abelian group generated by the zero-dimensional things in $\mathbb{S}^2$. Intuition tells us that the zero-dimensional things in $\mathbb{S}^2$ are precisely the individual points $x \in \mathbb{S}^2$, so we take $\text{FAb}(x \in \mathbb{S}^2)$ as an initial building block for $H_0(\mathbb{S}^2)$. Since all individual points are closed, the subgroup of zero-cycles is all of $\text{FAb}(x \in \mathbb{S}^2)$, too. In order to arrive at $H_0(\mathbb{S}^2)$, we have to take the quotient of $\text{FAb}(x \in \mathbb{S}^2)$ with respect to the homology equivalence relation.

⁸The free Abelian group is free because it has no relations other than those required by the Abelian group axioms.

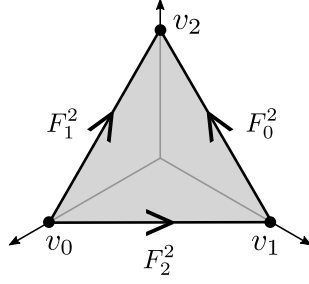


Figure 2.4: The standard two-simplex $\Delta^2 = [v_0, v_1, v_2]$ in $\mathbb{R}^3$. The interior $\mathring{\Delta}^2$ is shaded in light grey and the vertices v_0, v_1, v_2 are shown as solid black circles. The faces F_0^2, F_1^2, F_2^2 making up the boundary $\partial\Delta^2$ are drawn in solid black colour with arrows indicating their induced orientation.

The latter determines that two points $x, y \in \mathbb{S}^2$ are homologous if they differ by the boundary ∂c of some one-dimensional curve c , so the homology equivalence class $[x]$ of any given point $x \in \mathbb{S}^2$ reads

$$[x] = \{y \in \mathbb{S}^2 \mid x - y = \partial c \text{ for some one-dimensional curve } c \subset \mathbb{S}^2\}. \quad (2.26)$$

Note that the condition $x - y = \partial c$ makes sense because within the framework of the free Abelian group the points in $\mathbb{S}^2$ are allowed to have integer coefficients. The equivalence class Eq. (2.26) of any point $x \in \mathbb{S}^2$ contains all points that can be connected to it by a one-dimensional curve within $\mathbb{S}^2$. But since $\mathbb{S}^2$ is path-connected, this means that after factoring out the homologous points there is only a single equivalence class with integer coefficients left and we get

$$H_0(\mathbb{S}^2) \simeq \mathbb{Z}. \quad (2.27)$$

More generally, we will find that the zeroth homology group $H_0(X)$ of a topological space X is equal to

$$H_0(X) = \mathbb{Z}^{N_c(X)}, \quad (2.28)$$

where $N_c(X)$ is the number of connected components of X .

Singular and Simplicial Homology

In the previous section, we presented a graphic construction of the zeroth homology group of the two-sphere. In doing so, we relied on our intuition for low-dimensional objects: it is easy to see that zero-dimensional objects without boundary are just points, and that the boundaries of one-dimensional objects are pairs of points. In order to generalise these ideas, we now need a formal framework that applies to higher-dimensional objects. Specifically, we need a systematic way to define n -cycles and n -boundaries in a topological space X . Both singular homology and simplicial homology rely on a geometric structure called n -simplices to achieve this.

Definition 2.1.17. An **n -simplex** is the convex hull of any generic $n + 1$ points in $\mathbb{R}^{n+1}$ where generic means that they do not sit in a subspace of smaller dimension. The **standard n -simplex** is

$$\Delta^n = \left\{ \sum_{i=0}^n t_i v_i \mid \sum_i t_i = 1, t_i \geq 0 \right\} := [v_0, \dots, v_n], \quad (2.29)$$

where the **vertices** $\{v_0, \dots, v_n\}$ of Δ^n are the canonical basis vectors of $\mathbb{R}^{n+1}$.

In more graphical terms, an n -simplex is a generalised n -dimensional triangle: a 0-simplex is a point, a 1-simplex is a line segment, a 2-simplex is a triangle, a 3-simplex is a tetrahedron, and so on. Importantly, there is a natural definition for the faces and the boundary of an n -simplex. Let Δ^n be an n -simplex, then the j -th face F_j^n of Δ^n is the subsimplex spanned by all vertices but the j -th one, i.e.

$$F_j^n = \left\{ \sum_{i \neq j}^n t_i v_i \mid \sum_{i \neq j} t_i = 1, t_i \geq 0 \right\}. \quad (2.30)$$

From this, we can define the boundary as the union $\partial\Delta^n = \cup_{j=0}^n F_j^n$ of all faces, and the interior as $\mathring{\Delta}^n = \Delta^n \setminus \partial\Delta^n$. Figure 2.4 shows how the above definitions apply to $\Delta^2 \subset \mathbb{R}^3$. Note that every standard simplex Δ^n is defined with respect to the standard ordering $[v_0, \dots, v_n]$ of its vertices. This determines an orientation for the simplex. This orientation is hereditary, which means that it induces an orientation on the faces of the original simplex, cf. the arrows in Fig. 2.4. The key idea of singular and simplicial homology is to use continuous maps

$$\sigma^n : \Delta^n \rightarrow X \quad (2.31)$$

from the standard n -simplex into a given topological space X to define the simple n -dimensional things⁹ in X and then exploit the natural notion of boundaries for n -simplices to arrive at a sensible definition of n -cycles and n -boundaries in X . In the preliminary considerations, we based the definition of $H_0(\mathbb{S}^2)$ on the free Abelian group of zero-dimensional things. Accordingly, the formal definition of the n -th singular and simplicial homology groups is based on the free Abelian group

$$\mathcal{C}_n(X) := \text{FAb}(\sigma_\alpha^n : \Delta^n \rightarrow X) \quad (2.32)$$

generated by a family $\sigma := \{\sigma_\alpha^n\}$ of n -simplices in X . The fundamental difference between singular and simplicial homology lies in the particular family of n -simplices that is being used. We will equip objects of singular homology with a subscript or superscript “s”, while we will furnish objects of simplicial homology with a subscript or superscript “ Δ ”. The reasons for this will soon become apparent.

Singular homology puts no restrictions on the maps in $\sigma_s^n := \{\sigma_{s,\alpha}^n\}$ – they must be continuous, but they need not be injective and there may even be non-equivalent simplices with the same image in X . For example, the constant map

$$\begin{aligned} \sigma_s^2 : \Delta^2 &\rightarrow X \\ s &\mapsto x_0 \end{aligned} \quad (2.33)$$

is always continuous, even though it is clearly not injective. This makes the constant map an example of a *singular* map. Note that Eq. (2.33) also has the same image as the map

$$\begin{aligned} \sigma_s^0 : \Delta^0 &\rightarrow X \\ s &\mapsto x_0, \end{aligned} \quad (2.34)$$

which is both continuous and injective. The fact that the n -simplices are allowed to be singular in this way is what gives singular homology its name.

In contrast, simplicial homology deploys a highly regularised set $\sigma_\Delta^n := \{\sigma_{\Delta,\alpha}^n\}$ of n -simplices. It originates from the notion of replacing a potentially complicated topological space X with a much simpler triangulated space that is homeomorphic to it. To find such an equivalent triangulated space, we use a structure called a Δ -complex.¹⁰ A Δ -complex structure on a topological space X can be thought of as a well-behaved partition of X into a simple *network* of simplices.

Definition 2.1.18. Let X be a topological space. A **Δ -complex structure** on X is a collection

$$\sigma_\Delta^n := \{\sigma_\alpha^n : \Delta^n \rightarrow X\} \quad (2.35)$$

of continuous attaching maps σ_α^n , such that

- a) the images of interiors cover X , i.e. each $\sigma_\alpha^n|_{\mathring{\Delta}^n}$, is injective and every point of X is in the image of exactly one of these restrictions,
- b) if F is a face of Δ^n then $\sigma_\alpha^n|_F$ is one of the maps $\sigma_\beta^{n-1} : \Delta^{n-1} \rightarrow X$ after identifying F with Δ^{n-1} by the canonical homomorphism that preserves the vertex ordering, and
- c) a subset $U \subseteq X$ is open if and only if $(\sigma_\alpha^n)^{-1}(U)$ is open for all σ_α^n .

For a given topological space X , the family σ_Δ^n of n -simplices in simplicial homology is then defined as a suitable Δ -complex structure on X . Note once more that not every space admits a Δ -complex structure. Those that do are called *triangularisable* and usually admit many distinct Δ -complex structures.

⁹These n -dimensional things, the images $\sigma^n(\Delta^n) \subseteq X$ of n -simplices in X , are often called the *n -simplices of X* .

¹⁰The “ Δ ” in Δ -complex is a pictogram of a simplex, indicating that Δ -complexes are constructed from simplices.

Let X be a topological space. The free Abelian group $\mathcal{C}_n(X)$ in Eq. (2.32) is called the n -th chain group of X and the elements of $\mathcal{C}_n(X)$ are generally called the n -chains of X . The n -cycles and n -boundaries form normal Abelian subgroups of $\mathcal{C}_n(X)$ that can be defined using the natural notion of a boundary on the level of n -simplices. Specifically, we introduce a boundary homomorphism $\partial_n : \mathcal{C}_n(X) \rightarrow \mathcal{C}_{n-1}(X)$ that generalises intuitive results like the boundary

$$\partial_1(v_0 \bullet \rightarrow \bullet v_1) = v_1 - v_0 \quad (2.36)$$

of a one-chain that we discussed in the preliminary considerations on the zeroth homology group of $\mathbb{S}^2$.

Definition 2.1.19. Let X be a topological space and let $\mathcal{C}_n(X) = \text{FAb}(\sigma^n)$ denote the n -th singular ($\sigma^n \equiv \sigma_s^n$) or simplicial ($\sigma^n \equiv \sigma_\Delta^n$) chain group. An n -chain $C \in \mathcal{C}_n(X)$ is by definition a finite formal sum $C = \sum_\alpha c_\alpha \sigma_\alpha^n$ with $c_\alpha \in \mathbb{Z}$ and $\sigma_\alpha^n \in \sigma^n$. We define the **n -th boundary homomorphism** $\partial_n : \mathcal{C}_n(X) \rightarrow \mathcal{C}_{n-1}(X)$ by specifying its values on the basis elements σ_α^n of $\mathcal{C}_n(X)$ as

$$\partial_n(\sigma_\alpha^n) := \sum_{i=0}^n (-1)^i \sigma_\alpha |_{[v_0, \dots, \widehat{v}_i, \dots, v_n]}, \quad (2.37)$$

where $\widehat{v}_i$ means that v_i is omitted in the domain of the attaching map σ_α^n .

Based on the boundary homomorphisms we can uniquely define the subgroups of n -cycles and n -boundaries in $\mathcal{C}_n(X)$. The n -cycles are defined as the elements $C \in \mathcal{C}_n(X)$ with zero boundary

$$\partial_n(C) = 0. \quad (2.38)$$

Thus, the n -cycles are precisely the kernel of the n -th boundary map. We define

$$Z_n(X) := \ker(\partial_n) = \{C \in \mathcal{C}_n(X) \mid \partial_n(C) = 0\}. \quad (2.39)$$

Since the kernel of any group homomorphism is a normal subgroup, the n -cycles form a normal subgroup $Z_n(X) \triangleleft \mathcal{C}_n(X)$. Furthermore, $Z_n(X)$ is Abelian because every subgroup of an Abelian group is Abelian. Similarly, the n -boundaries are defined as the images of the $(n+1)$ -th boundary map, i.e.

$$B_n(X) := \text{im}(\partial_{n+1}) = \{\partial_{n+1}(C) \mid C \in \mathcal{C}_{n+1}(X)\}. \quad (2.40)$$

Since the image of any group homomorphism is a subgroup, and every subgroup of an Abelian group is Abelian and normal, the n -boundaries form a normal Abelian subgroup $B_n(X) \triangleleft \mathcal{C}_n(X)$. In order to see that $B_n(X)$ is also a normal subgroup of $Z_n(X)$, i.e. $B_n(X) \triangleleft Z_n(X)$, we note that the boundary homomorphisms between chain groups satisfy the essential property

$$\partial_n \circ \partial_{n+1} = 0, \quad (2.41)$$

which is often stated as $\partial^2 = 0$ in the literature. Crucially, Eq. (2.41) tells us that the image of ∂_{n+1} is contained in the kernel of ∂_n , i.e.

$$\text{im}(\partial_{n+1}) \subseteq \ker(\partial_n), \quad (2.42)$$

or, equivalently, that $B_n(X) \subseteq Z_n(X)$. This allows us to consider the $(n+1)$ -th boundary map ∂_{n+1} as a homomorphism

$$\partial_{n+1} : \mathcal{C}_{n+1}(X) \rightarrow Z_n(X), \quad (2.43)$$

which immediately shows us that $B_n(X) \triangleleft Z_n(X)$ by the same reasoning as before. Since $B_n(X)$ is a normal subgroup of $Z_n(X)$, the definition

$$H_n(X) = Z_n(X)/B_n(X) = \ker(\partial_n)/\text{im}(\partial_{n+1}) \quad (2.44)$$

of the n -th singular or simplicial homology group of X is always well-defined and always results in another Abelian group. It is worth pausing to appreciate that this is why we specifically use Abelian groups for this construction. In Abelian groups, every subgroup is normal by definition so every subgroup gives rise to a well-defined quotient group. Furthermore, all subgroups, quotients and direct sums of Abelian groups are again Abelian. A given topological space X usually accommodates simplices of many different orders so it makes sense to come up with a language that allows for a comprehensive study of the homology of X . This is done in terms of chain complexes.

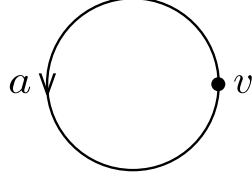


Figure 2.5: The minimal Δ -complex structure of $\mathbb{S}^1$ consisting of a single vertex v and a single edge a that connects v to itself.

Definition 2.1.20. A **chain complex** $(A_\bullet, d_\bullet)$ is a sequence $A_\bullet = \{A_i\}_{i \in I}$ of Abelian **chain groups** A_i connected by a family $d_\bullet = \{d_i : A_i \rightarrow A_{i-1}\}$ of homomorphisms d_i called **boundary maps** such that the composition of any two consecutive maps always is the zero map, i.e. $d_i \circ d_{i-1} = 0$ for all $i \in I$. We write

$$\dots \xrightarrow{d_{n+2}} A_{n+1} \xrightarrow{d_{n+1}} A_n \xrightarrow{d_n} A_{n-1} \xrightarrow{d_{n-1}} \dots \quad (2.45)$$

and define its **chain homology** as

$$H_n(A_\bullet, d_\bullet) = \ker d_n / \operatorname{im} d_{n+1} . \quad (2.46)$$

For a given topological space X , the sequence of singular or simplicial chain groups $\mathcal{C}_\bullet = \{\mathcal{C}_i\}_{i \in I}$ together with the boundary homomorphisms $\partial_\bullet = \{\partial_i : \mathcal{C}_i \rightarrow \mathcal{C}_{i-1}\}$ forms a singular or simplicial n -chain complex $(\mathcal{C}_\bullet, \partial_\bullet)$ with a well-defined chain homology. We can therefore define singular and simplicial homology of X as follows.

Definition 2.1.21. Let X be any topological space and let $(\mathcal{C}_\bullet^s, \partial_\bullet)$ be its singular n -chain complex. The n -th **singular homology group** $H_n^s(X)$ is defined as the chain homology

$$H_n^s(X) := H_n(\mathcal{C}_\bullet^s(X), \partial_\bullet) \quad (2.47)$$

of the singular chain complex.

Definition 2.1.22. Let X be a triangularisable topological space and let $(\mathcal{C}_\bullet^\Delta, \partial_\bullet)$ be its simplicial n -chain complex. The n -th **simplicial homology group** $H_n^\Delta(X)$ is defined as the chain homology

$$H_n^\Delta(X) := H_n(\mathcal{C}_\bullet^\Delta(X), \partial_\bullet) \quad (2.48)$$

of the simplicial chain complex.

Note once more that the unregulated nature of singular homology makes it applicable to any topological space X while the definition of simplicial homology requires X to be triangularisable. In a way, this makes singular homology a superior homology theory. However, for triangularisable spaces X the singular homology groups and the simplicial homology groups are equivalent. Simplicial homology may therefore be understood as a tool for computing singular homology for triangularisable spaces. In the following, we explicitly compute the simplicial homology groups of two important low-dimensional manifolds.

Simplicial Homology: Examples

Earlier we determined the fundamental group $\pi_1(\mathbb{S}^1)$ of the circle $\mathbb{S}^1$ to get an idea of homotopy groups. Let us return to this example and determine the simplicial homology of $X = \mathbb{S}^1$. We can represent $\mathbb{S}^1$ by a minimal Δ -complex structure that consists of a single vertex v and a single edge a connecting v to itself, cf. Fig. 2.5. This results in the simplicial chain complex

$$\begin{array}{ccccccc} 0 & \xrightarrow{\iota} & \mathcal{C}_1(\mathbb{S}^1) & \xrightarrow{\partial_1} & \mathcal{C}_0(\mathbb{S}^1) & \xrightarrow{P_0} & 0 , \\ & & \downarrow \wr & & \downarrow \wr & & \\ & & \mathbb{Z} & & \mathbb{Z} & & \end{array} \quad (2.49)$$

where we added trivial maps $\partial_2 =: \iota$ and $\partial_0 =: P_0$ from and to the trivial group $0 \equiv \{0\}$ because the definitions of the zeroth and first homology groups require knowledge of $\text{im}(\partial_2)$ and $\text{ker}(\partial_0)$. Specifically, the leftmost homomorphism $\partial_2 =: \iota$ is the trivial inclusion map of 0 into the highest non-trivial chain group $\mathcal{C}_1(\mathbb{S}^1)$ of the simplicial chain complex, and the map $\partial_0 =: P_0$ is the trivial projection map from lowest non-trivial chain group $\mathcal{C}_0(\mathbb{S}^1)$ of the simplicial chain complex to the trivial group. The only potentially non-trivial homomorphism in the simplicial chain complex is the boundary map ∂_1 , which maps any one-simplex to the difference between its end vertex and its beginning vertex. However, our Δ -complex structure of $\mathbb{S}^1$ only features a single one-simplex and a single vertex. Accordingly, we have $\partial_1 a = v - v = 0$, so that ∂_1 is trivial as well, i.e. $\partial_1 = 0$. Both $\mathcal{C}_1(\mathbb{S}^1) = \text{FAb}(\sigma^1 : \Delta^1 \rightarrow a)$ and $\mathcal{C}_0(\mathbb{S}^1) = \text{FAb}(\sigma^0 : \Delta^0 \rightarrow v)$ are free Abelian groups on a single object so they are both isomorphic to the integers as $\mathcal{C}_0(\mathbb{S}^1) \simeq \mathcal{C}_1(\mathbb{S}^1) \simeq \mathbb{Z}$. With this we compute

$$H_0^\Delta(\mathbb{S}^1) = \text{ker}(\partial_0) / \text{im}(\partial_1) = \text{ker}(P_0) / \text{im}(\partial_1) . \quad (2.50)$$

Now the kernel of P_0 is all of $\mathcal{C}_0(\mathbb{S}^1) \simeq \mathbb{Z}$ by definition. The image of $\partial_1 = 0$ is of course $\{0\} \in \mathcal{C}_0(\mathbb{S}^1) \simeq \mathbb{Z}$, so we get

$$H_0^\Delta(\mathbb{S}^1) = \mathbb{Z}/0 \simeq \mathbb{Z} . \quad (2.51)$$

Note that although the quotient of the integers by 0 might seem alarming at first glance, it is perfectly well-defined. This is because the quotient is taken in the sense of cosets of the normal subgroup $\{0\} \triangleleft \mathbb{Z}$, which simply recovers $\mathbb{Z}$ itself. Analogously, we compute

$$H_1^\Delta(\mathbb{S}^1) = \text{ker}(\partial_1) / \text{im}(\partial_2) = \text{ker}(\partial_1) / \text{im}(\iota) . \quad (2.52)$$

Again, the kernel of $\partial_1 = 0$ trivially gives all of $\mathcal{C}_0(\mathbb{S}^1) \simeq \mathbb{Z}$ while the image of $\partial_2 = \iota$ is equal to $\{0\} \in \mathcal{C}_0(\mathbb{S}^1) \simeq \mathbb{Z}$ by definition, and we find

$$H_1^\Delta(\mathbb{S}^1) = \mathbb{Z}/0 \simeq \mathbb{Z} . \quad (2.53)$$

Combined, the homology groups of $\mathbb{S}^1$ are

$$H_n^\Delta(\mathbb{S}^1) = \text{ker } \partial_n / \text{im } \partial_{n+1} = \begin{cases} \mathbb{Z} & n = 0, 1 \\ 0 & \text{else} . \end{cases} \quad (2.54)$$

This result is in line with our initial considerations regarding the zeroth homology group and its relation to the number of connected components. It also reproduces the earlier result for the fundamental group of the circle, namely

$$\pi_1(\mathbb{S}^1) \simeq H_1^\Delta(\mathbb{S}^1) \simeq \mathbb{Z} . \quad (2.55)$$

In fact, the latter is a special case of a more general statement, namely that the first homology group is always equal to (an Abelianisation of) the fundamental group if X is connected. Yet, our analysis of the simplicial homology of $\mathbb{S}^1$ has provided us with some additional insights, too. In particular, it tells us that $H_0(\mathbb{S}^1)$ and $H_1(\mathbb{S}^1)$ are the *only* non-trivial simplicial homology groups of $\mathbb{S}^1$, so there are no strange higher-dimensional holes that could be detected by higher simplicial homology groups. Importantly, the result in Eq. (2.54) can be generalised to all higher-dimensional spheres $\mathbb{S}^d$ as

$$H_n^\Delta(\mathbb{S}^d) = \begin{cases} \mathbb{Z} & n = 0, d \\ 0 & \text{else} . \end{cases} \quad (2.56)$$

The homology of the d -sphere plays an important role in defining orientability for topological manifolds. The concept of orientability will become important later on, so we briefly outline its general idea here. Let M be a closed, connected n -dimensional manifold. By definition, every point $x \in M$ has an open neighbourhood $U_x \subset M$ that is homeomorphic to $\mathbb{R}^n$. To study the topological structure around x , we focus on how x interacts with the surrounding space. This is done by removing x from M as $M \setminus \{x\}$. After removing x , the neighbourhood $U_x \setminus \{x\}$ becomes homeomorphic to $\mathbb{R}^n \setminus \{\mathbf{0}\}$, which is in turn

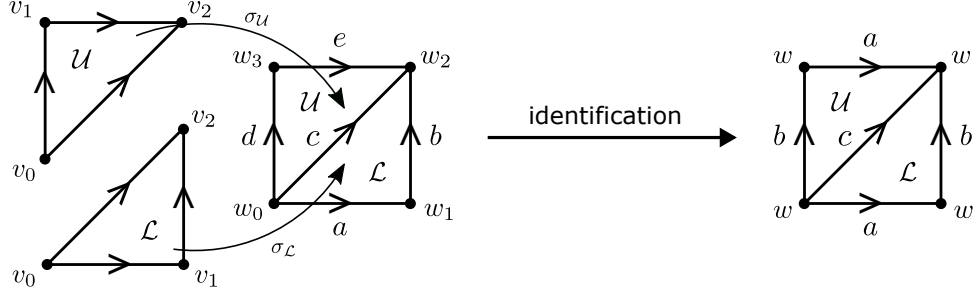


Figure 2.6: The minimal Δ -complex structure of $\mathbb{T}^2$ consisting of two separate 2-simplices $\mathcal{U}$ and $\mathcal{L}$ (leftmost picture) that are glued together along one edge to give five 1-simplices a, b, c, d, e and four vertices w_0, w_1, w_2, w_3 (right part of left picture). Upon identification, the edge a is glued to e , the edge b is glued to d and all vertices are glued together. The resulting Δ -complex structure of $\mathbb{T}^2$ features two 2-simplices $\mathcal{U}$ and $\mathcal{L}$, three distinct edges a, b, c and a single vertex w (rightmost picture).

homotopic to the $(n-1)$ -sphere $\mathbb{S}^{n-1}$. In this sense, the local structure of M near x is characterised¹¹ by $\mathbb{S}^{n-1}$. The homology of $\mathbb{S}^{n-1}$ then serves as a tool to define a local orientation of M at x as a choice of generator $\mu_x \in \{-1, +1\}$ of $H_{n-1}(\mathbb{S}^{n-1}) \simeq \mathbb{Z}$. If it is possible to consistently choose the same generator for all $x \in M$, we say that M is orientable and an orientation of M is defined as a function $x \mapsto \mu_x$ that assigns a local orientation to all $x \in M$ in a consistent way. If a closed, connected n -dimensional manifold M is orientable, its top homology group $H_n(M)$ is isomorphic to the integers, i.e. $H_n(M) \simeq \mathbb{Z}$, and the n -th homology class that represents the chosen global generator of $H_n(M)$ is called the *fundamental class* or *orientation class* $[M]$ of M .

The one-sphere is an especially accessible textbook example, which is further distinguished by the fact that it has a certain relevance for physical theories. It appears, for instance, as the one-dimensional Brillouin zone in the theory of condensed matter. Another simple topological manifold that has a similar significance for physical theories is the manifold that serves as a two-dimensional Brillouin zone: the two-torus $\mathbb{T}^2$. We will discuss its homology as a final example in the following.

The basis for the minimal Δ -complex structure of $\mathbb{T}^2$ consists of two two-dimensional patches $\mathcal{U}$ and $\mathcal{L}$, five one-dimensional edges a, b, c, d, e and four zero-dimensional vertices w_0, w_1, w_2, w_3 . In order to turn this Δ -complex structure into a space that is homeomorphic to $\mathbb{T}^2$ we need to identify $w \sim w_0 \sim w_1 \sim w_2 \sim w_3$, $a \sim e$ and $b \sim d$. As a consequence, the original building blocks of the Δ -complex structure reduce to two distinct two-dimensional patches $\mathcal{U}$ and $\mathcal{L}$, three distinct one-dimensional edges a, b, c and only one zero-dimensional vertex w , cf. Fig. 2.6. The non-trivial n -chain groups $\mathcal{C}_n(\mathbb{T}^2)$ are therefore the free Abelian groups

$$\begin{aligned} \mathcal{C}_0(\mathbb{T}^2) &= \text{FAb}(\sigma^0 : \sigma^0 \rightarrow w) \simeq \mathbb{Z}, \\ \mathcal{C}_1(\mathbb{T}^2) &= \text{FAb}(\sigma_1^1 : \sigma^1 \rightarrow a, \sigma_2^1 : \sigma^1 \rightarrow b, \sigma_3^1 : \sigma^1 \rightarrow c) \simeq \mathbb{Z}^3, \\ \mathcal{C}_2(\mathbb{T}^2) &= \text{FAb}(\sigma_1^2 : \sigma^2 \rightarrow \mathcal{U}, \sigma_2^2 : \sigma^2 \rightarrow \mathcal{L}) \simeq \mathbb{Z}^2. \end{aligned} \quad (2.57)$$

and the associated simplicial chain complex is

$$\begin{array}{ccccccc} 0 & \xrightarrow{\iota} & \mathcal{C}_2(\mathbb{T}^2) & \xrightarrow{\partial_2} & \mathcal{C}_1(\mathbb{T}^2) & \xrightarrow{\partial_1} & \mathcal{C}_0(\mathbb{T}^2) \xrightarrow{P_0} 0 \\ & & \wr & & \wr & & \wr \\ & & \mathbb{Z}^2 & & \mathbb{Z}^3 & & \mathbb{Z} \end{array} \quad (2.58)$$

with the same meaning of symbols as before. In order to determine the simplicial homology groups we have to sort out the boundary maps. Once more there is only a single zero-simplex so we have $\partial_1 = 0$ with

$$\ker(\partial_1) = \mathcal{C}_1(\mathbb{T}^2) \simeq \mathbb{Z}^3 \quad \text{and} \quad \text{im}(\partial_1) = \{0\}. \quad (2.59)$$

¹¹The link between the local structure and the homology of $\mathbb{S}^{n-1}$ is provided by the concept of relative singular homology, which goes beyond the scope of this thesis but can, for instance, be found in Ref. [15].

For ∂_2 we have to work a little bit harder. First, we explicitly compute the boundary maps of $\mathcal{U}$ and $\mathcal{L}$ separately as

$$\partial_2 \mathcal{U} = +\sigma_{\mathcal{U}}|_{[v_1, v_2]} - \sigma_{\mathcal{U}}|_{[v_0, v_2]} + \sigma_{\mathcal{U}}|_{[v_0, v_1]} = +a - c + b \quad (2.60)$$

and

$$\partial_2 \mathcal{L} = +\sigma_{\mathcal{L}}|_{[v_1, v_2]} - \sigma_{\mathcal{L}}|_{[v_0, v_2]} + \sigma_{\mathcal{L}}|_{[v_0, v_1]} = +a - c + b, \quad (2.61)$$

where we followed the leftmost and the rightmost pictures in Fig. 2.6. This may seem a bit counter-intuitive at first because usually the “positive” mathematical direction is counter-clockwise. Going counter-clockwise around the $\mathcal{L}$ patch on the right of Fig. 2.6 and adding up the boundary edges with their orientation sign produces $b - c + a$ just like in Eq. (2.61). However, repeating the same procedure for patch $\mathcal{U}$ yields $-a - b + c$, i.e. the opposite result of the one in Eq. (2.60). So why does assuming the standard positive “counter-clockwise” orientation give the wrong result here? The reason for this is that the original $\mathcal{U}$ two-simplex is not in its standard orientation. This can be seen as follows. Going counter-clockwise around the individual $\mathcal{L}$ on the left of Fig. (2.6) gives $v_0 \rightarrow v_1 \rightarrow v_2$, which corresponds to the standard vertex order that defines the orientation of the $\mathcal{L}$ patch. Doing the same for the individual $\mathcal{U}$ yields $v_0 \rightarrow v_2 \rightarrow v_1$, which is an odd permutation of the standard vertex order producing the opposite orientation for the $\mathcal{U}$ patch. Note that reversing the orientation of either $\mathcal{U}$ or $\mathcal{L}$ in Fig. (2.6) is required to unambiguously glue them together along the c edge, which connects $w_0 = v_0$ to $w_2 = v_2$. The fact that the boundaries of $\mathcal{U}$ and $\mathcal{L}$ are the same, $\partial_2 \mathcal{U} = \partial_2 \mathcal{L}$, tells us that

$$\ker \partial_2 = \{j\mathcal{U} + k\mathcal{L} \mid j + k = 0\} \simeq \mathbb{Z}, \quad (2.62)$$

because the boundary map is linear. The middle expression is isomorphic to the integers $\mathbb{Z}$ because the constraint $j + k = 0$ removes one degree of freedom and thus reduces $\mathbb{Z}^2 \simeq \{j\mathcal{U} + k\mathcal{L} \mid j, k \in \mathbb{Z}\}$ to $\mathbb{Z}$. The image of ∂_2 is simply the free Abelian group of a single generator, namely

$$\text{im } \partial_2 = \text{FAb}(a + b - c) = \mathbb{Z}. \quad (2.63)$$

The three simplicial homology groups of $\mathbb{T}^2$ are therefore

$$\begin{aligned} H_2^\Delta(\mathbb{T}^2) &= \ker \partial_2 / \text{im } \iota = \mathbb{Z}/0 = \mathbb{Z}, \\ H_1^\Delta(\mathbb{T}^2) &= \ker \partial_1 / \text{im } \partial_2 = \mathbb{Z}^3 / \mathbb{Z} = \mathbb{Z}^2, \\ H_0^\Delta(\mathbb{T}^2) &= \ker P_0 / \text{im } \partial_1 = \mathbb{Z}/0 = \mathbb{Z}. \end{aligned} \quad (2.64)$$

The two-torus is path-connected, which is readily confirmed by $H_0^\Delta(\mathbb{T}^2) \simeq \mathbb{Z}$. Furthermore, the first simplicial homology group $H_1^\Delta(\mathbb{T}^2) \simeq \mathbb{Z}^2$ matches the fundamental group $\pi_1(\mathbb{T}^2) \simeq \mathbb{Z}^2$ that we computed earlier. As was mentioned before, the first homology group of a connected topological space X is always equal to the Abelianisation of its fundamental group, so this is expected as well. Finally, $H_2^\Delta(\mathbb{T}^2) \simeq \mathbb{Z}$ signifies that $\mathbb{T}^2$ is an orientable topological manifold, which is also a well-established fact. Note that the n -th Betti number β_n is really defined as the rank

$$\beta_n := \text{rank}(H_n(X)) \quad (2.65)$$

of the n -th *homology* group, rather than the n -th *homotopy* group as suggested in Eq. (2.20). The definition of the first Betti number via the fundamental group that we gave in Eq. (2.20) poses a very delicate special case: it holds only because $\text{rank}(\pi_1(M)) = \text{rank}(H_1(M))$ for two-dimensional closed orientable manifolds M [65]. Usually, the Abelianisation process relating $\pi_1(X)$ to $H_1(X)$ eliminates generators and reduces the rank. One example of this is the *non-orientable* Klein bottle $\mathbb{K}^2$, where $\text{rank}(\pi_1(\mathbb{K}^2)) = 2$, but $\text{rank}(H_1(\mathbb{K}^2)) = 1$.

Overall, we find that the topological information provided by the homology groups corresponds well to our natural notion of the respective manifolds. Still, it is natural to ask how much of simplicial homology depends on the chosen Δ -complex structure on X as opposed to X itself. Above, we used the minimal Δ -complex structures to compute the simplicial homology of $\mathbb{S}^1$ and $\mathbb{T}^2$. Of course, one can always add

more simplices to the description. For example, we could use a Δ -complex structure for $\mathbb{S}^1$ that consists of two vertices v_0 and v_1 connected by two edges a and b such that the initial non-trivial n -chain groups would be $\mathcal{C}_0(\mathbb{S}^1) \simeq \mathcal{C}_1(\mathbb{S}^1) \simeq \mathbb{Z}^2$ resulting in a simplicial chain complex

$$0 \xrightarrow{\iota} \mathcal{C}_1(\mathbb{S}^1) \xrightarrow{\partial_1} \mathcal{C}_0(\mathbb{S}^1) \xrightarrow{P_0} 0. \quad (2.66)$$

$\begin{array}{ccc} \wr & & \wr \\ \mathbb{Z}^2 & & \mathbb{Z}^2 \end{array}$

However, the two vertices v_0 and v_1 in the new Δ -complex structure mean that the boundary map ∂_1 is no longer trivial. Instead, its image and kernel can be shown to be isomorphic to the integers, i.e. $\ker(\partial_1) \simeq \text{im}(\partial_1) \simeq \mathbb{Z}$, such that we still get

$$\begin{aligned} H_0^\Delta(\mathbb{S}^1) &= \ker(P_0)/\text{im}(\partial_1) = \mathbb{Z}^2/\mathbb{Z} \simeq \mathbb{Z} \\ H_1^\Delta(\mathbb{S}^1) &= \ker(\partial_1)/\text{im}(\iota) = \mathbb{Z}/0 \simeq \mathbb{Z}. \end{aligned} \quad (2.67)$$

Indeed, singular homology can be used to show that the simplicial homology theory of a triangularisable topological space X does not depend on the choice of Δ -complex structure.

There is a natural generalisation of the homology theory we have discussed so far that can offer certain technical advantages in some cases. It involves a simple modification of the underlying simplicial chain complexes. Up to this point, the n -chain groups $\mathcal{C}_n(X)$ were taken to be free Abelian groups on the n -simplices, i.e. their elements were finite formal sums $\sum_\alpha g_\alpha \sigma_\alpha^n$ with integer coefficients $g_\alpha \in \mathbb{Z}$. Now, there is no compelling reason to restrict homology theory to integer coefficients. The generalisation occurs when we allow coefficients from any Abelian group G so the n -chain groups become the Abelian groups

$$\mathcal{C}_n(X; G) := \left\{ \sum_\alpha g_\alpha \sigma_\alpha^n \mid \sigma_\alpha^n : \Delta^n \rightarrow X, g_\alpha \in G \right\} \quad (2.68)$$

of finite formal G -sums of n -simplices in X . All subsequent machinery carries over so we can define the homology theory with coefficients in the same way as before. The resulting homology groups with coefficients in G are denoted by $H_n(X; G)$. Common coefficient groups include $\mathbb{Z}_p$, $\mathbb{Q}$, and $\mathbb{R}$. There is an important theorem called the universal coefficient theorem that tells us that the integral homology groups $H_n(X) = H_n(X; \mathbb{Z})$ completely determine the homology groups $H_n(X; G)$ with coefficients. Nonetheless, some notions depend strongly on the coefficient group. Orientability, for example, becomes redundant with $\mathbb{Z}_2$ coefficients because the two possible distinct generators $\{-1, +1\}$ of $\mathbb{Z}$ are identified in $\mathbb{Z}_2$. As a consequence, *every* topological manifold is $\mathbb{Z}_2$ -orientable. For this reason, orientability is often noted as G -orientability, highlighting the respective coefficient group for clarity. If G is not explicitly mentioned, one is usually referring to $\mathbb{Z}$ -orientability.

2.1.10 Cohomology

As its name suggests, cohomology is a theory that results from a *dualisation* of homology. Accordingly, the cohomology groups $H^n(X)$ of a topological space X fulfil very similar axioms and capture largely the same topological information as the homology groups $H_n(X)$. In fact, the homology groups of a space fully determine its cohomology groups. The reason why we are still interested in cohomology is that it comes with an extra layer of algebraic structure, namely a natural product map $H^n(X) \times H^m(X) \rightarrow H^{n+m}(X)$, that is extremely useful in many situations. Furthermore, there are various topological contexts where cohomology arises naturally. One of these is the theory of characteristic classes of fibre bundles that is very important for topological condensed matter theory. As with homology, there are different cohomology theories. In particular, there are the directly dualised versions of simplicial and singular homology that are known as simplicial and singular cohomology, respectively. For this part, we closely follow Ref. [61].

We have mentioned that cohomology results from homology through dualisation. So what does this mean? Duality in mathematics is less a fixed concept than a general principle, a recurring theme that manifests across nearly every branch of the field. As such, it lacks a single unambiguous definition. Yet, at its core, duality always describes a deep correspondence between mathematical objects that appear unrelated, but somehow encode the same underlying structure – as if viewed from different perspectives. Familiar examples include the duality between open and closed sets in topological spaces

X or that between a finite-dimensional inner product space $(V, \langle \cdot, \cdot \rangle)$ over a field F and its dual vector space $V^* := \text{hom}(V, F)$ of linear maps from V to F . The concrete duality operation that connects the open and closed sets is the operation of taking the complement, while the concrete duality operation between vectors and linear functions in the inner product space is the operation of applying the canonical isomorphism that maps $v \mapsto \langle v, \cdot \rangle$. To understand the duality operation that takes us from the homology of a topological space to its cohomology, we first discuss how the algebraic dualisation the underlying chain-complex works. To this end, consider a general chain-complex $(A_\bullet, d_\bullet)$ of free Abelian groups A_n connected by boundary maps d_n as

$$\dots \longrightarrow A_{n+1} \xrightarrow{d_{n+1}} A_n \xrightarrow{d_n} A_{n-1} \longrightarrow \dots \quad (2.69)$$

To dualise this chain-complex we first replace each chain group A_n by its dual group, the cochain group

$$A_n^*(R) = \text{hom}(A_n, R), \quad (2.70)$$

where R denotes any ring¹² of coefficients. Common coefficient rings include $R = \mathbb{Z}_p, \mathbb{Q}, \mathbb{R}$, but a standard choice for R is again the integers $R = \mathbb{Z}$. Unless otherwise stated, we will use $R = \mathbb{Z}$ and no longer mention R explicitly in the following. Next, we replace the boundary map $d_n : A_n \rightarrow A_{n-1}$ by its dual map

$$d_n^* : A_{n-1}^* \rightarrow A_n^*, \quad (2.71)$$

which we call the coboundary map. Note that the coboundary map d_n^* is a homomorphism from A_{n-1}^* to A_n^* rather than a homomorphism from A_n^* to A_{n-1}^* . This is a direct consequence of the definition of a dual homomorphism: given any homomorphism $\varphi : A \rightarrow B$ its dual homomorphism φ^* is defined by composition

$$\begin{aligned} \varphi^* : \text{hom}(B, R) &\rightarrow \text{hom}(A, R) \\ \psi &\mapsto \varphi^*(\psi) := \psi \circ \varphi. \end{aligned} \quad (2.72)$$

As a consequence, φ^* maps a function $\psi : B \rightarrow R$ to a new function $\psi \circ \varphi : A \rightarrow R$, which effectively reverses the direction of φ . For the boundary map, this implies that $d_n^* : A_{n-1}^* \rightarrow A_n^*$ is now labelled by the codomain A_n^* rather than the domain A_{n-1}^* . In order to simplify general statements about homology and cohomology we define the dual map

$$\delta_n := d_{n+1}^*, \quad (2.73)$$

that is in line with the notation we know from homology, i.e. δ_n is a homomorphism from A_n^* to A_{n+1}^* labelled by its domain. With the coboundary maps δ_n the cochain groups A_n^* form a chain-complex

$$\dots \longleftarrow A_{n+1}^* \xleftarrow{\delta_n} A_n^* \xleftarrow{\delta_{n-1}} A_{n-1}^* \longleftarrow \dots \quad (2.74)$$

which is of the same form as the original chain complex in Eq. (2.69) and that is called the cochain complex $(A_\bullet^*, \delta_\bullet)$ of $(A_\bullet, d_\bullet)$. Duality ensures that the important property of $\partial^2 = 0$ is directly inherited as $\delta^2 = 0$ by the coboundary map. The chain cohomology of $(A_\bullet, d_\bullet)$ can therefore be defined as the chain homology of the cochain complex.

Definition 2.1.23. Let $(A_\bullet, d_\bullet)$ be a chain complex of free Abelian groups and let $(A_\bullet^*, \delta_\bullet)$ denote its cochain complex. The **chain cohomology** of $(A_\bullet, d_\bullet)$ is defined as the chain homology of the cochain complex $(A_\bullet^*, \delta_\bullet)$, i.e. the n -th chain cohomology group $H^n(A)$ is defined as

$$H^n(A_\bullet, d_\bullet) = \ker \delta_n / \text{im } \delta_{n-1}. \quad (2.75)$$

With this, the singular and simplicial cohomology groups $H_s^n(X)$ and $H_\Delta^n(X)$ of a topological space X can be defined as follows.

¹²Technically the coefficient ring could be relaxed to an Abelian coefficient group, but the most common choices for R are rings so we use this slightly specialised definition.

Definition 2.1.24. Let X be any topological space with a singular n -chain complex $(\mathcal{C}_\bullet^s, \partial_\bullet)$ and let $H_n^s(X) := H_n(\mathcal{C}_\bullet^s(X), \partial_\bullet)$ denote its singular homology groups. The **singular cohomology** groups $H_s^n(X)$ of X are defined as the chain cohomology

$$H_s^n(X) := H^n(\mathcal{C}_\bullet^s(X), \partial_\bullet) = \ker \delta_n / \text{im } \delta_{n-1} . \quad (2.76)$$

Definition 2.1.25. Let X be any triangularisable topological space with a simplicial n -chain complex $(\mathcal{C}_\bullet^\Delta, \partial_\bullet)$ and let $H_n^\Delta(X) := H_n(\mathcal{C}_\bullet^\Delta(X), \partial_\bullet)$ denote its simplicial homology groups. The **simplicial cohomology** groups $H_\Delta^n(X)$ of X are defined as the chain cohomology

$$H_\Delta^n(X) := H^n(\mathcal{C}_\bullet^\Delta(X), \partial_\bullet) = \ker \delta_n / \text{im } \delta_{n-1} . \quad (2.77)$$

Just as in homology theory, singular cohomology can be used to show that the simplicial cohomology does not depend on the specific choice of Δ -complex structure of a triangularisable topological space.

In the following, we will discuss some general properties of cohomology that do not depend on the specific cohomology type. For this reason, we will omit sub- and superscripts referencing the cohomology type whenever possible. The duality of cohomology and homology is captured by the pairing map

$$\begin{aligned} H^n(X) \times H_n(X) &\rightarrow R \\ ([\alpha], [\sigma]) &\mapsto [\alpha]([\sigma]) , \end{aligned} \quad (2.78)$$

that assigns a scalar from the coefficient ring R to any pair $([\alpha], [\sigma])$ of the n -th cohomology and homology classes. Importantly, the cohomology groups are completely determined by the coefficient ring R and the homology groups. Despite this, cohomology has an additional piece of algebraic structure that homology lacks. Specifically, there exists a natural bilinear map

$$\begin{aligned} \smile : H^n(X) \times H^m(X) &\rightarrow H^{n+m}(X) \\ ([\alpha], [\beta]) &\mapsto [\alpha] \smile [\beta] , \end{aligned} \quad (2.79)$$

called the cup product, which is directly induced by a cup product of cochains. There are some subtleties in the definition of the cup product that occur when the coefficient ring R is only an Abelian G group but we will not discuss such cases here [15].

Definition 2.1.26. Let $\mathcal{C}^n(X)$ and $\mathcal{C}^m(X)$ denote the n -th and m -th cochain group of a topological space X . For $\alpha \in \mathcal{C}^n(X)$ and $\beta \in \mathcal{C}^m(X)$ we define the **cup product** $\alpha \smile \beta \in \mathcal{C}^{n+m}(X)$ as the cochain whose value on an $(n+m)$ -simplex $\sigma : \Delta^{n+m} = [v_0, \dots, v_{n+m}] \rightarrow X$ is given by

$$(\alpha \smile \beta)(\sigma) = \alpha(\sigma|_{[v_0, \dots, v_n]}) \cdot \beta(\sigma|_{[v_n, \dots, v_{n+m}]}) , \quad (2.80)$$

where the $\cdot$ denotes the multiplication of the coefficient ring.

In order to understand how this cup product of cochains induces a cup product of cohomology classes, we have to relate it to the coboundary map δ_n . The idea is to define $\delta_{n+m}(\alpha \smile \beta)$ such that

$$\delta_{n+m}(\alpha \smile \beta)(\sigma) \stackrel{!}{=} (\alpha \smile \beta)(\partial_{n+m}\sigma) \quad (2.81)$$

for every $(n+m)$ -simplex σ . One can show that this is precisely the case when

$$\delta_{n+m}(\alpha \smile \beta) = \delta_n \alpha \smile \beta + (-1)^n \alpha \smile \delta_m \beta . \quad (2.82)$$

From Eq. (2.82) it is immediately clear that the cup product $\alpha \smile \beta$ of two cocycles is again a cocycle, i.e. that $\delta_{n+m}(\alpha \smile \beta) = 0$ for every α and β with $\delta_n \alpha = \delta_m \beta = 0$. Furthermore, the cup product of a cocycle α and a coboundary $\delta_m \beta$ is again a coboundary because $\delta_{n+m}(\alpha \smile \beta) = (-1)^n \alpha \smile \delta_m \beta$ where we used that $\delta_n \alpha = 0$. Consequently, there is an induced cup product Eq. (2.79) that is associative and distributive. The latter allows us to promote the direct sum

$$H^*(X) := \bigoplus_n H^n(X) \quad (2.83)$$

to a graded ring $(H^*(X), \smile)$ by endowing it with the cup product. The ring structure of the cohomology groups under the cup product has important implications for the analysis of topological spaces. For example, it can sometimes be used to show that two spaces with isomorphic (co)homology groups are different because they have different cohomology rings. Another application is that the cohomology ring can be used to generate rich algebraic invariants. This is crucial for the field of characteristic classes and characteristic numbers of bundle spaces.

Cohomology and homology theories can be generalised using an axiomatic approach: the so-called Eilenberg–Steenrod axioms, specify a list of properties that all (co)homology theories must have in common. Singular (co)homology is the prime example of a theory that satisfies these axioms, making it the prototypical theory for defining and comparing new (co)homology theories. However, there are other important (co)homology theories that extend or specialise singular (co)homology, often to adapt it to a particular class of topological spaces.

de Rham Cohomology

There are many different ways to define cohomology theories for topological spaces. Examples include simplicial and singular cohomology that we discussed earlier, but also various other cohomology theories that are tailored to specific classes of topological spaces. For some exotic spaces, the different cohomology theories yield different answers. However, there is also a large class of topological spaces on which they all agree. Within this class of spaces, some of the more specialised cohomology theories can be understood as tools for computing the singular cohomology for a particular type of space. For example, simplicial cohomology can be used to compute singular cohomology on triangularisable topological spaces. Another very relevant example that we will discuss in the next section is the theory of *de Rham cohomology* which allows the computation of singular cohomology for smooth manifolds: if a topological space M is a smooth manifold, its cohomology can be naturally defined in terms of differential forms on M . The following part is mostly based on Ref. [39].

Let us start by refreshing some of the general definitions about differential forms. Let V be an n -dimensional real vector space. A type (p, k) tensor is a multi-linear map

$$T : \bigoplus^p V^* \bigoplus^k V \rightarrow \mathbb{R}$$

$$(\omega^1, \dots, \omega^p, v_1, \dots, v_k) \mapsto T(\omega^1, \dots, \omega^p, v_1, \dots, v_k), \quad (2.84)$$

that maps p dual vectors $\omega^1, \dots, \omega^p$ and k vectors $v_1, \dots, v_k$ to a real number $T(\omega^1, \dots, \omega^p, v_1, \dots, v_k)$. The integers p and k indicate the number of contravariant and covariant indices, respectively. This terminology refers to the behaviour of the tensor under a basis transformation A : while the vectors $v_i \in V$ transform *covariantly*, via A itself, the dual vectors $\omega^i \in V^*$ transform *contravariantly*, via the inverse transformation A^{-1} . The definition of differential forms builds on the notion of covariant k -tensors, i.e. tensors of type $(0, k)$. Specifically, it requires the concept of *alternating* covariant k -tensors φ , for which a permutation of the arguments multiplies the value by the sign of the permutation, i.e.

$$\varphi(v_{\pi(1)}, \dots, v_{\pi(k)}) = \text{sign}(\pi) \varphi(v_1, \dots, v_k). \quad (2.85)$$

Here, $\pi \in S_k$ is a permutation in the k -th symmetric group and the sign of π is defined as

$$\text{sign}(\pi) = (-1)^{N_\pi} \quad (2.86)$$

with the minimal number N_π of transpositions that make up π . Alternating covariant k -tensors are called exterior k -forms and the space of all exterior k -forms on V is denoted by $\bigwedge^k(V^*)$. For $\dim(V) = n$ the dimension of $\bigwedge^k(V^*)$ is $\dim(\bigwedge^k(V^*)) = \binom{n}{k}$ so the direct sum

$$\bigwedge(V^*) = \bigoplus_{k=0}^n \bigwedge^k(V^*) \quad (2.87)$$

is a vector space of dimension $\dim(\bigwedge(V^*)) = \sum_{k=0}^n \binom{n}{k} = 2^n$. Now, the space $\bigwedge(V^*)$ becomes an algebra

when equipped with the so-called wedge (or exterior) product

$$\begin{aligned} \wedge : \bigwedge^k(V^*) \times \bigwedge^l(V^*) &\rightarrow \bigwedge^{k+l}(V^*) \\ (\varphi, \psi) &\mapsto \varphi \wedge \psi, \end{aligned} \quad (2.88)$$

where

$$(\varphi \wedge \psi)(v_1, \dots, v_{k+l}) := \sum_{\pi \in S_{k+l}} \text{sign}(\pi) \varphi(v_{\pi(1)}, \dots, v_{\pi(k)}) \psi(v_{\pi(k+1)}, \dots, v_{\pi(k+l)}) \quad (2.89)$$

is the antisymmetrised tensor product of φ and ψ . The wedge product is bilinear, associative and anticommutates as

$$\varphi \wedge \psi = (-1)^{kl} \psi \wedge \varphi. \quad (2.90)$$

Let M be an n -dimensional smooth manifold. As we will soon discuss in greater detail, the tangent bundle TM of M is based on the disjoint collection of tangent spaces at all points $p \in M$ and smooth vector fields X on M are defined as smooth maps $X : M \rightarrow TM$. We may repeat this construction for exterior k -forms to get the exterior k -form bundle

$$\bigwedge^k(T^*M) = \bigsqcup_{p \in M} \bigwedge^k(T_p^*M), \quad (2.91)$$

which can be used to define differential forms by analogy to the definition of vector fields.

Definition 2.1.27. Let M denote a smooth manifold. A continuous section

$$\omega : M \rightarrow \bigwedge^k(T^*M) \quad (2.92)$$

of the exterior k -form bundle $\bigwedge^k(T^*M)$ is called a **k -form** on M . The integer k is called the degree of the k -form.

By definition, a k -form ω is a tensor field that assigns an alternating tensor to every point of M . The vector space of all k -forms on a smooth manifold M is denoted $\Omega^k(M)$ and the wedge product of two k -forms is defined pointwise, i.e.

$$(\omega \wedge \eta)_p := \omega_p \wedge \eta_p, \quad (2.93)$$

such that the wedge product of a k -form and an l -form is a $(k+l)$ -form. The direct sum

$$\Omega(M) := \bigoplus_{k=0}^n \Omega^k(M) \quad (2.94)$$

of the $\Omega^k(M)$ is a vector space which also forms an associative anticommutative graded algebra $(\Omega(M), \wedge)$ with the wedge product. In any smooth chart (U, φ) of a smooth manifold M , a k -form $\omega \in \Omega^k(M)$ can locally be written as

$$\omega = \sum_{i_1, \dots, i_k} \omega_{i_1, \dots, i_k}(x^1, \dots, x^n) dx^{i_1} \wedge \dots \wedge dx^{i_k}, \quad (2.95)$$

where the coefficients $\omega_{i_1, \dots, i_k}(x^1, \dots, x^n)$ are smooth functions on the $(x^1, \dots, x^n)$ -coordinate patch. One can show that for each smooth manifold M there exists a differential operator [39]

$$d_k : \Omega^k(M) \rightarrow \Omega^{k+1}(M), \quad (2.96)$$

called the exterior derivative, that satisfies

$$d^2 := d_{k+1} \circ d_k = 0, \quad (2.97)$$

and is given by

$$d\omega = \sum_{i_1, \dots, i_k} d\omega_{i_1, \dots, i_k} \wedge dx^{i_1} \wedge \dots \wedge dx^{i_k} = \sum_{j, i_1, \dots, i_k} (\partial_j \omega_{i_1, \dots, i_k}) dx^j \wedge dx^{i_1} \wedge \dots \wedge dx^{i_k}. \quad (2.98)$$

on any chart of M . The exterior derivative allows us to define the closed and exact differential forms.

Definition 2.1.28. Let M be an n -dimensional smooth manifold. A differential k -form $\omega \in \Omega^k(M)$ on M is called **closed** if $d\omega = 0$ and **exact** if there exists a $(k-1)$ -form $\eta \in \Omega^{k-1}(M)$ such that $\omega = d\eta$.

The set of closed k -forms constitutes a group called the k -th cocycle group $Z^k(M)$ and the set of exact k -forms constitutes a group called the k -th coboundary group $B^k(M)$. Due to Eq. (2.97), we find

$$B^k(M) \triangleleft Z^k(M) . \quad (2.99)$$

The de Rham cohomology groups $H_{\text{dR}}^n(M)$ of an n -dimensional smooth manifold M can therefore be defined as the homology of the cochain complex

$$\dots \longleftarrow \Omega^{k+1}(M) \xleftarrow{d_k} \Omega^k(M) \xleftarrow{d_{k-1}} \Omega^{k-1}(M) \longleftarrow \dots \quad (2.100)$$

Definition 2.1.29. Let M be an n -dimensional smooth manifold. The **k -th de Rham cohomology group** $H_{\text{dR}}^k(M)$ of M is defined as

$$H_{\text{dR}}^k(M) = \ker(d_k) / \text{im}(d_{k-1}) = Z^k(M) / B^k(M) . \quad (2.101)$$

Note that the de Rham cohomology groups are naturally defined with real coefficients because the underlying alternating covariant k -tensors are maps into the real numbers by definition. For every $\omega \in Z^k(M)$ the equivalence class $[\omega] \in H_{\text{dR}}^k(M)$ is defined as

$$[\omega] = \{ \eta \in Z^k(M) \mid \eta = \omega + d\alpha \text{ for } \alpha \in \Omega^{k-1}(M) \} , \quad (2.102)$$

i.e. two forms are equivalent, or *cohomologous*, if they differ by an exact form. Just like singular cohomology, de Rham cohomology is equipped with a cup product: it is directly induced by the wedge product in Eq. (2.88) as

$$\begin{aligned} \wedge : H_{\text{dR}}^n(X) \times H_{\text{dR}}^m(X) &\rightarrow H_{\text{dR}}^{m+n}(X) \\ ([\alpha], [\beta]) &\mapsto [\alpha] \wedge [\beta] := [\alpha \wedge \beta] . \end{aligned} \quad (2.103)$$

The above construction shows that de Rham cohomology is a well-defined cohomology theory. However, it is not immediately clear which homology theory it is dual to. It turns out that de Rham cohomology is a dualisation of singular homology. This dual relationship is captured by the natural de Rham pairing

$$\begin{aligned} \langle \cdot, \cdot \rangle : \mathcal{C}_k(M) \times \Omega^k(M) &\rightarrow \mathbb{R} \\ (\alpha, \omega) &\mapsto \langle \alpha, \omega \rangle := \int_{\alpha} \omega . \end{aligned} \quad (2.104)$$

between singular k -chains $\alpha \in \mathcal{C}_k^s(M)$ and de Rham k -cochains $\omega \in \Omega^k(M)$ of an n -dimensional smooth manifold M . In another important step, one can use Eq. (2.104) to show that de Rham cohomology is in fact isomorphic to singular cohomology with real coefficients. To do this, we define a family $\phi = \{\phi_k\}_{k=1, \dots, n}$ of homomorphisms

$$\begin{aligned} \phi_k : \Omega^k(M) &\rightarrow \mathcal{C}_k^*(M) \\ \omega &\mapsto \phi_k(\omega) := \langle \cdot, \omega \rangle , \end{aligned} \quad (2.105)$$

which take de Rham k -cochains $\omega \in \Omega^k(M)$ to singular k -cochains $\phi_k(\omega) \in \mathcal{C}_k^*(M)$ of M . The idea is that the ϕ_k naturally induce homomorphisms

$$\begin{aligned} \Phi_k : H_{\text{dR}}^k(M) &\rightarrow H_s^k(M; \mathbb{R}) \\ [\omega] &\mapsto \Phi_k([\omega]) := [\phi_k(\omega)] \end{aligned} \quad (2.106)$$

between the k -th de Rham cohomology group and the k -th singular cohomology group with real coefficients. However, this is not trivially the case. For one thing, we have to make sure that the Φ_k are well-defined. Furthermore, the ϕ_k from Eq. (2.105) have to survive the cohomology construction in

Def. (2.1.29) so they must take cycles to cycles and boundaries to boundaries. The striking insight by de Rham was that all of this is guaranteed by Stokes theorem

$$\int_{\partial\alpha} \omega = \int_{\alpha} d\omega. \quad (2.107)$$

This can be seen as follows. First, we note that we can use Eq. (2.104) to rewrite Stokes theorem as

$$\langle \partial\alpha, \omega \rangle = \langle \alpha, d\omega \rangle, \quad (2.108)$$

which is readily compatible with the inner product Eq. (2.104) that gives rise to the ϕ_k in Eq. (2.105). In order to show that Eq. (2.106) is well-defined we have to prove two things: the well-definedness of the maps $\Phi_k : H_{\text{dR}}^k(M) \rightarrow H_s^k(M; \mathbb{R})$ and the well-definedness of the $\Phi_k([\omega]) \in H_s^k(M)$ as maps $\Phi_k([\omega]) : H_k^s(M) \rightarrow \mathbb{R}$ themselves. The latter means that

$$(\Phi_k([\omega]))([\alpha]) := (\phi_k(\omega))(\alpha) \quad (2.109)$$

does not depend on the representative α of $[\alpha] \in H_k^s(M)$. To show this, we take two representatives α and α' of the equivalence class $[\alpha]$. By definition $[\alpha]$ contains cycles that are homologous, i.e. all k -cycles α with $\partial\alpha = 0$ that differ by a k -boundary $\beta = \partial\xi$ for a suitable $(k+1)$ -chain ξ . This means we can write

$$\alpha' = \alpha + \partial\xi. \quad (2.110)$$

If we plug this into the right-hand side of Eq. (2.109) we get

$$\begin{aligned} (\phi_k(\omega))(\alpha') &= (\phi_k(\omega))(\alpha + \partial\xi) \\ &= \langle \alpha + \partial\xi, \omega \rangle \\ &\stackrel{(\diamond)}{=} \langle \alpha, \omega \rangle + \langle \partial\xi, \omega \rangle \\ &\stackrel{(*)}{=} \langle \alpha, \omega \rangle + \langle \xi, d\omega \rangle \\ &\stackrel{(*)}{=} \langle \alpha, \omega \rangle \\ &= (\phi_k(\omega))(\alpha), \end{aligned} \quad (2.111)$$

where we used the bilinearity of $\langle \cdot, \cdot \rangle$ in $(\diamond)$, Stokes theorem Eq. (2.108) in $(*)$, and the fact that ω is a cocycle, i.e. that $d\omega = 0$, in $(*)$. The calculation in Eq. (2.111) shows that Eq. (2.109) is independent of the choice of representative $\alpha \in [\alpha]$ and therefore well-defined. The well-definedness of $\Phi_k : H_{\text{dR}}^k(M) \rightarrow H_s^k(M; \mathbb{R})$ works analogously but this time we take two representatives ω and ω' of the equivalence class $[\omega]$. By definition $[\omega]$ contains cocycles that are cohomologous, i.e. all k -forms ω with $d\omega = 0$ that differ by an exact form $\epsilon = d\eta$. This means we can write

$$\omega' = \omega + d\eta \quad (2.112)$$

for a suitable $(k-1)$ -form η and repeat the calculation we did in Eq. (2.111). Next we have to show that the ϕ_k are “compatible” with the boundary maps $d_\bullet$ and $\delta_\bullet$ of the de Rham cochain complex $(\Omega^\bullet(M), d_\bullet)$ and the singular cochain complex $(\mathcal{C}_\bullet^*(M; \mathbb{R}), \delta_\bullet)$. Mathematically, we say that ϕ is compatible with the boundary maps if it is a chain map, i.e. if it is a family of homomorphisms that commutes with the boundary maps $\partial_\bullet$ and $d_\bullet$ of the chain complexes. This is the case if the diagram

$$\begin{array}{ccccccc} \dots & \xleftarrow{d_{k+1}} & \Omega^{k+1}(M) & \xleftarrow{d_k} & \Omega^k(M) & \xleftarrow{d_{k-1}} & \dots \\ & & \downarrow \phi_{k+1} & & \downarrow \phi_k & & \\ \dots & \xleftarrow{\delta_{k+1}} & \mathcal{C}_{k+1}^*(M; \mathbb{R}) & \xleftarrow{\delta_k} & \mathcal{C}_k^*(M; \mathbb{R}) & \xleftarrow{\delta_{k-1}} & \dots \end{array} \quad (2.113)$$

commutes, i.e. if

$$\phi_{k+1} \circ d_k = \delta_k \circ \phi_k \quad (2.114)$$

for every k in the chain complexes. Since Eq. (2.114) is a rather abstract relation between two maps

$$\phi_{k+1} \circ d_k : \Omega^k(M) \rightarrow \mathcal{C}_{k+1}^*(M) \quad \text{and} \quad \delta_k \circ \phi_k : \Omega^k(M) \rightarrow \mathcal{C}_{k+1}^*(M), \quad (2.115)$$

we plug in a test k -cochain $\omega^k \in \Omega^k(M)$ to get

$$(\phi_{k+1} \circ d_k)(\omega^k) \in \mathcal{C}_{k+1}^*(M) \quad (2.116)$$

and a test $(k+1)$ -chain $\alpha_{k+1} \in \mathcal{C}_{k+1}(M)$ to get

$$((\phi_{k+1} \circ d_k)(\omega^k))(\alpha_{k+1}) \in \mathbb{R}. \quad (2.117)$$

With this we can write

$$\begin{aligned} ((\phi_{k+1} \circ d_k)(\omega^k))(\alpha_{k+1}) &= (\phi_{k+1}(d_k \omega^k))(\alpha_{k+1}) \\ &= \langle \alpha_{k+1}, d_k \omega^k \rangle \\ &\stackrel{(\diamond)}{=} \langle \partial_{k+1} \alpha_{k+1}, \omega^k \rangle \\ &= (\phi_k(\omega^k))(\partial_{k+1} \alpha_{k+1}) \\ &= (\phi_k(\omega^k) \circ \partial_{k+1})(\alpha_{k+1}) \\ &\stackrel{(\star)}{=} (\partial_{k+1}^* (\phi_k(\omega^k)))(\alpha_{k+1}) \\ &\stackrel{(*)}{=} (\delta_k (\phi_k(\omega^k)))(\alpha_{k+1}) \\ &= ((\delta_k \circ \phi_k)(\omega^k))(\alpha_{k+1}), \end{aligned} \quad (2.118)$$

where we used Stokes theorem from Eq. (2.108) in $(\diamond)$, the definition Eq. (2.72) of the dual homomorphism in $(\star)$, and the definition $\delta_k := d_{k+1}^*$ of the coboundary map from Eq. (2.73) in $(*)$. The calculation Eq. (2.118) proves that the diagram in Eq. (2.113) commutes. Accordingly, ϕ is a chain map and induces a family of cohomology homomorphisms as in Eq. (2.106). The de Rham theorem then asserts that Eq. (2.106) is an isomorphism for every smooth manifold M . We formally state it for the sake of completeness [39].

Theorem 2.1.2. de Rham Theorem. *For every smooth manifold M the map*

$$\begin{aligned} \Phi_k : H_{\text{dR}}^k(M) &\rightarrow H_s^k(M; \mathbb{R}) \\ [\omega] &\mapsto [\langle \cdot, \omega \rangle] \end{aligned} \quad (2.119)$$

is a well-defined isomorphism that establishes the equivalence of de Rham cohomology and singular cohomology with real coefficients.

The de Rham theorem tells us that we can use the rather tangible machinery of de Rham cohomology, i.e. differential forms and integration on manifolds, to compute the singular cohomology with real coefficients for smooth manifolds. This makes de Rham cohomology a very powerful tool. One of its applications, which is particularly relevant for physics, is its appearance in the theory of characteristic classes.

2.2 Fibre Bundles

An n -dimensional manifold is a topological space that locally looks like $\mathbb{R}^n$. Similarly, a fibre bundle is a topological space that locally looks like a direct product of two topological spaces. The interesting thing about both manifolds and fibre bundles is that their global structure is usually more complicated than their simple local structure would suggest. Therefore, answers to questions will generally differ depending on whether they are asked at a local or global level: while local questions can be answered based on the simpler structures of a Euclidean space or a product space, answering global questions requires the complete information of the manifold or fibre bundle, respectively. This indicates that manifolds and fibre bundles are particularly relevant when it comes to “global” questions. The following section is mostly based on Refs. [39] and [15].

The importance of manifolds for theoretical physics has been appreciated at least since the introduction of general relativity, where they provide a natural framework for the description of curved spacetime [39]. The significance of fibre bundles is a little bit more subtle, but very closely related. It shows when we ask follow-up questions like “How do electromagnetic fields spread over spacetime?” The inclusion of the electromagnetic fields means that we require a mathematical structure that takes into account not only the spacetime manifold itself, but also the electromagnetic fields on it: at each point in space and time we have to store the local geometric data of the electromagnetic field vectors. Fibre bundles formalise the attachment of local geometric data (fibres) to an underlying manifold (base) and thus provide the natural mathematical framework for all theories of this general type.

One prominent class of such theories are gauge theories which are concerned with the organisation of a Lie group, called the gauge group, over an underlying manifold [39]. The preceding example of electromagnetism is a special case of a gauge theory. Other examples are Lagrangian systems where the phase space is realised as the tangent bundle of the configuration space [67], as well as classical and quantum field theories where fields are vector-valued and operator-valued sections over spacetime, respectively [39]. Formally, these theories cover much of modern theoretical physics. Since fibre bundles capture the global properties of many physical theories, and topology is generally concerned with global properties, fibre bundles also control many of the topological phenomena in physics. These include the tenfold way of topological insulators and superconductors [14], instantons in gauge theories [7, 30–32], vortices in superfluids [36–38], and skyrmions in both quantum field theory and condensed matter theory [4, 5, 33]. As was mentioned in the introduction of this thesis, instantons, monopoles, vortices and skyrmions can be understood as special cases of a more general class of objects called topological solitons, which ultimately correspond to non-perturbative, localised and topologically stable solutions to field equations [30]. For this reason, the relevance of fibre bundles in physics is sometimes associated with the non-perturbative aspects of (quantum) field theories. Since several of these topological phenomena will appear later, we briefly review the fundamentals of fibre bundle theory below. Let us begin with a definition [39].

Definition 2.2.1. Fibre Bundle. A fibre bundle is a structure (E, B, F, π) consisting of

1. a topological space E called the **total space**,
2. a topological space B called the **base space**,
3. a topological space F called the **typical fibre**, and
4. a surjective continuous map $\pi : E \rightarrow B$ called the **projection map**,

satisfying the following condition: for every $p \in B$ there exists an open neighbourhood $U \subset B$ of p with a homeomorphism $\phi : U \times F \rightarrow \pi^{-1}(U)$ that is compatible with the projection map as $(\pi \circ \phi)(p, f) = p$. Such a map ϕ is called a **local trivialisation** of the bundle over U .

The existence of a local trivialisation around every $p \in B$ means that the preimage $\pi^{-1}(p) =: F_p$ of the projection map over every $p \in B$ is homeomorphic to the typical fibre F as $F_p \simeq F$. For this reason, it is appropriate to picture a fibre bundle, at least locally, as a brush or a strand of hair. The important thing is that these local strands can be twisted and warped as we go around the base space, thus breaking with this simple picture on the global scale.

It is common practice to denote a fibre bundle (E, B, F, π) by $E \xrightarrow{\pi} B$ or simply E for brevity. The latter can serve as a gentle reminder that, despite the considerable complexity of Def. 2.2.1, a fibre bundle is ultimately just a topological space E that can, but need not, be described in the rich way of a fibre bundle. The above definition solely relies on topological concepts like topological spaces, continuous functions, and homeomorphisms – no differentiability or explicit choice of coordinates is required for the construction of the bundle. Of course, one can modify Def. 2.2.1 in these regards. If we require all topological spaces and maps of to be differentiable or smooth we call the resulting fibre bundle a differentiable or smooth fibre bundle. Similarly, we can choose a specific atlas $\{(U_i, \phi_i)\}$, i.e. a collection of local trivialisation charts (U_i, ϕ_i) whose open neighbourhoods U_i cover B . The resulting bundle $(E, B, F, \pi, \{U_i\}, \{\phi_i\})$ is called a coordinate bundle. A fibre bundle as per Def. 2.2.1 is then an equivalence class of coordinate bundles under a simple equivalence relation where two coordinate bundles $(E, B, F, \pi, \{U_i\}, \{\phi_i\})$ and $(E, B, F, \pi, \{V_i\}, \{\psi_i\})$ are equivalent if $(E, B, F, \pi, \{U_i\} \cup \{V_i\}, \{\phi_i\} \cup \{\psi_i\})$ is again a coordinate bundle. In physics, we usually work with smooth coordinate bundles.

Consider a smooth coordinate bundle $(E, B, F, \pi, \{U_i\}, \{\phi_i\})$. It is natural to wonder how the transitions between overlapping trivialisation charts (U_i, ϕ_i) and (U_j, ϕ_j) work. After all, only these transitions can account for any non-trivial twisting of the bundle. To examine this, we consider a trivialisation chart (U_i, ϕ_i) and use the local trivialisation $\phi_i : U_i \times F \rightarrow \pi^{-1}(U_i)$ to define a map

$$\begin{aligned} \phi_{i,p}(f) : F &\rightarrow F \\ f &\mapsto \phi_{i,p}(f) := \phi_i(p, f), \end{aligned} \quad (2.120)$$

which is a diffeomorphism for every $p \in U_i \subset B$. If two trivialisation charts (U_i, ϕ_i) and (U_j, ϕ_j) have a non-trivial overlap $U_i \cap U_j \neq \emptyset$ we demand that

$$\phi_{j,p}(f) \stackrel{!}{=} \phi_{i,p}(t_{ij,p}(f)) \quad (2.121)$$

for all $p \in (U_i \cap U_j)$. Here, the $t_{ij,p}$ are smooth maps

$$\begin{aligned} t_{ij,p} : F &\rightarrow F \\ f &\mapsto t_{ij,p}(f) := (\phi_{i,p}^{-1} \circ \phi_{j,p})(f), \end{aligned} \quad (2.122)$$

called the transition functions between the charts (U_i, ϕ_i) and (U_j, ϕ_j) . The transition functions of a fibre bundle are often required to live in some group G of symmetry transformations, which determines the matching conditions between overlapping trivialisation charts. In such cases, the group G is required to act continuously and faithfully on the fibre F on the left, which means that there exists a group action

$$\begin{aligned} \rho : G &\rightarrow \text{Aut}(F) \\ g &\mapsto \rho(g) \end{aligned} \quad (2.123)$$

with $\ker(\rho) = \{e\}$, such that the map

$$\begin{aligned} \Phi : G \times F &\rightarrow F \\ (g, f) &\mapsto \Phi(g, f) := \rho(g)(f) \equiv g \cdot f \end{aligned} \quad (2.124)$$

is continuous.¹³ Here, $\text{Aut}(F)$ is the group of all diffeomorphisms from F to itself and the trivial kernel of ρ means that ρ is injective and hence faithful. The continuity of Φ defines the continuity of the group action and we often use the symbolic notation $g \cdot f$ of left multiplication to describe the left action of $g \in G$ on $f \in F$. In fibre bundles with a symmetry group G it is then required that the transition functions $t_{ij,p} : F \rightarrow F$ correspond to a representation $\rho(g) : F \rightarrow F$ for some $g \in G$ as

$$t_{ij,p} = \rho(g), \quad (2.125)$$

¹³Note that the continuity of Φ is only well-defined when G is understood as a topological group, i.e. a group that is equipped with a topology that is *compatible* with the group operation in the sense that the group operation map $\cdot : G \times G \rightarrow G$, $(x, y) \mapsto x \cdot y$ and the inversion map $^{-1} : G \rightarrow G$, $x \mapsto x^{-1}$ are made continuous. Every group can be made into a topological group by considering it in the discrete topology. Lie groups are important topological groups with the standard manifold topology inherited from metric topology of $\mathbb{R}^n$ via the charts.

which is often casually stated as $t_{ij,p} \stackrel{!}{\in} G$. The group G is then called the structure group of the fibre bundle. In physics, we typically use smooth coordinate bundles with a Lie group G for a structure group and with smooth manifolds for the topological spaces E , B and F .

Definition 2.2.2. Physics Fibre Bundle. A **smooth coordinate G -bundle** or **(physics) fibre bundle** is a structure $(E, B, F, \pi, \{U_i\}, \{\phi_i\}, G)$ consisting of

1. a smooth manifold E called the **total space**,
2. a smooth manifold B called the **base space**,
3. a smooth manifold F called the **typical fibre**,
4. a surjective smooth map $\pi : E \rightarrow B$ called the **projection map**,
5. a collection of open charts $\{U_i\}$ that cover the base space B together with a collection $\{\phi_i\}$ of diffeomorphisms $\phi_i : U_i \times F \rightarrow \pi^{-1}(U_i)$ that establish the **coordinates** and the **local trivialisation** of the bundle,
6. a Lie group G called the **structure group**, that acts on F on the left and that supplies the transition functions $t_{ij,p} : F \rightarrow F$ for all $p \in (U_i \cap U_j) \neq \emptyset$ between overlapping trivialisation charts.

A fibre bundle of this kind is also often called a smooth fibre bundle to stress its differentiable structure. Importantly, a smooth fibre bundle as per Def. 2.2.2 is *itself* a smooth manifold so it possesses a dimension $\dim_{\mathbb{R}}(E)$ which is related to the dimension $\dim_{\mathbb{R}}(B)$ of the base manifold and the dimension $\dim_{\mathbb{R}}(F)$ of the fibre manifold via $\dim_{\mathbb{R}}(E) = \dim_{\mathbb{R}}(B) + \dim_{\mathbb{R}}(F)$. Let us consider some simple examples of topological spaces that form smooth fibre bundles as defined in Def. 2.2.2.

Example 2.2.1. Let $\mathbb{S}^1$ denote the unit one-sphere and let $I = (0, 1)$ denote the open unit interval. The following topological spaces are fibre bundles that can be constructed using $\mathbb{S}^1$ and I .

1. The cylinder C is a product space $C = \mathbb{S}^1 \times I$ and therefore a trivial fibre bundle with base space $B = \mathbb{S}^1$ and fibre $F = I$.
2. The Möbius strip M is a non-trivial fibre bundle with base space $B = \mathbb{S}^1$ and fibre $F = I$. Even though M looks like a product $I \times I$ locally, it exhibits a twist that distinguishes it from the trivial cylinder globally. The Möbius strip is a non-orientable topological space.
3. The two-torus $\mathbb{T}^2$ is a product space $C = \mathbb{S}^1 \times \mathbb{S}^1$ and therefore a trivial fibre bundle with base space $B = \mathbb{S}^1$ and fibre $F = \mathbb{S}^1$. It can be obtained from the cylinder by glueing together the boundary $\partial I = \{0\} \cup \{1\}$ of I .
4. The Klein bottle $\mathbb{K}^2$ is a non-trivial fibre bundle with base space $B = \mathbb{S}^1$ and fibre $F = \mathbb{S}^1$. Even though $\mathbb{K}^2$ looks like a product $\mathbb{S}^1 \times \mathbb{S}^1$ locally, it exhibits a twist that distinguishes it from the trivial two-torus globally. The Klein bottle is a non-orientable topological space.

Figure 2.7 shows three-dimensional instantiations of the four fibre bundles in Ex. 2.2.1. In the sketch of the Klein bottle $\mathbb{K}^2$, the size of the fibre $F = \mathbb{S}^1$ varies along the base manifold. There is no profound reason for this; it is simply the established strategy for illustrating the twist of the Klein bottle in three dimensions. In particular, the modification of the tube circumference has no topological significance as it constitutes a continuous deformation. Another feature of the standard depiction of $\mathbb{K}^2$ in Fig. 2.7 is its self-intersection. Importantly, this self-intersection is not an intrinsic property of $\mathbb{K}^2$ as a fibre bundle, but rather an artefact of the explicit instantiation of $\mathbb{K}^2$ in $\mathbb{R}^3$. It can be resolved in $\mathbb{R}^4$ much like the situation with the two interlinked rings we encountered in Fig. 2.1. The immersion of $\mathbb{K}^2$ in $\mathbb{R}^3$ is instructive nonetheless. It clearly shows that $\mathbb{K}^2$ is closed and non-orientable. The latter can be recognised by the fact that $\mathbb{K}^2$, like the Möbius strip M , has only one surface instead of the “inner” and “outer” surfaces of orientable manifolds like the two-torus $\mathbb{T}^2$ or the cylinder C . Another thing shown in Fig. 2.7 are sections of the cylinder and the Möbius bundle. These are represented as red lines. The formal definition of sections is as follows.

Definition 2.2.3. Sections of Fibre Bundles. Let $E \xrightarrow{\pi} B$ be a fibre bundle. A **(smooth) section** σ of $E \xrightarrow{\pi} B$ is a continuous (smooth) map $\sigma : B \rightarrow E$ which satisfies $\pi \circ \sigma = \text{id}_B$.

The condition $\pi \circ \sigma = \text{id}_B$ tells us that $\sigma(p)$ is an element of the fibre $F_p = \pi^{-1}(p)$ for every $p \in B$. The idea of a section is therefore to attach a *particular* element of the fibre, rather than the entire fibre, to every point of the base manifold. Importantly, this must be done in a continuous (smooth) fashion, which effectively means that the specific choices of elements over neighbouring points in B may only differ by

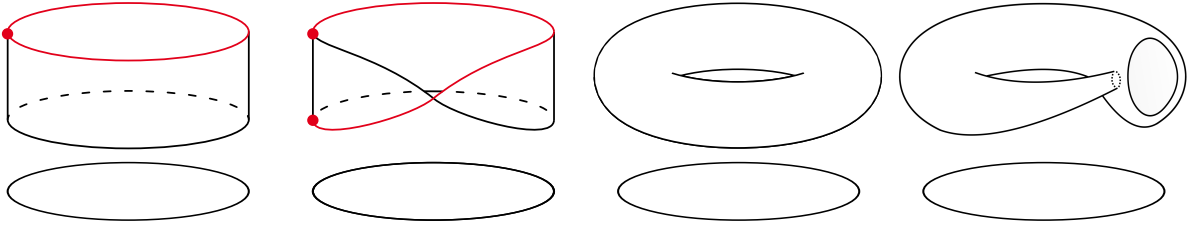


Figure 2.7: Three-dimensional sketches of four simple fibre bundles over $\mathbb{S}^1$. From left to right: a cylinder C with fibre $F = (0, 1)$, a Möbius strip M with fibre $F = (0, 1)$, a torus $\mathbb{T}^2$ with fibre $F = \mathbb{S}^1$, and a Klein bottle $\mathbb{K}^2$ with fibre $F = \mathbb{S}^1$. The four identical bottom circles portray the shared base manifold $B = \mathbb{S}^1$, the red lines in the pictures of C and M indicate sections (see text), and the dotted circle in the picture of $\mathbb{K}^2$ marks the self-intersection of the Klein bottle in $\mathbb{R}^3$. The three-dimensional instantiations of C , M , and $\mathbb{T}^2$ are embeddings into $\mathbb{R}^3$, while that of $\mathbb{K}^2$ is only an immersion. Illustration created by the author, inspired by Ref. [68].

a small amount. In this sense, a section is like a selected hairstyle on a head of (infinitely) long hair. If there exists a section σ that is well-defined for all of B at the same time, we call it a *global section*. The set of all global sections of a fibre bundle $E \xrightarrow{\pi} B$ is often denoted by $\Gamma(E, B)$. However, fibre bundles do not in general admit (non-trivial) global sections. The non-existence of (non-trivial) global sections is a result of the twistedness of the bundle and deeply rooted in topology. In fact, the obstruction to the existence of (non-trivial) global sections can often be measured by certain cohomology classes of the base space. The study of these obstruction cohomology classes is part of the theory of characteristic classes in algebraic topology. One particularly well-known example is at the heart of the hairy-ball theorem, where a characteristic class known as the Euler class obstructs the existence of a nowhere vanishing section of the tangent bundle over $\mathbb{S}^2$. We will go into more detail about characteristic classes and obstructions later on. Since not every fibre bundle admits (non-trivial) global sections, it makes sense to define local sections as continuous maps $\sigma_U : U \rightarrow E$ for open subsets $U \subset B$. The local triviality of fibre bundles ensures that local sections always exist. This makes them crucial tools in the study of fibre bundles.

There are some important special cases of fibre bundles that are characterised by the nature of the typical fibre. The most relevant ones for our purposes are called *vector bundles* and *principal bundles*.

2.2.1 Vector Bundles

Let E be a fibre bundle. We call E a vector bundle if its typical fibre F is a vector space over a field K . The dimension $\dim_K(F)$ of the fibre is then called the rank or sometimes dimension of the vector bundle and the structure group G is necessarily $G = \text{GL}(m, K)$. Typically, we encounter $K = \mathbb{R}, \mathbb{C}$ in physics.

Example 2.2.2. The tangent bundle TM over an m -dimensional manifold M is a vector bundle with base space $B = M$ and fibre $F = \mathbb{R}^m$. Using the more detailed tuple notation from Def. 2.2.2 we can therefore write TM as a vector bundle $(E, B, F, \pi, G) = (TM, M, \mathbb{R}^m, \pi, \text{GL}(m, \mathbb{R}))$ where we omitted the choice of trivialisation atlas. The global sections $\Gamma(TM, M)$ of TM are the vector fields on M .

A vector bundle whose fibre is a one-dimensional vector space over a field K is called a K -line bundle or simply a line bundle. The cylinder and the Möbius strip from Ex. 2.2.1 can be realised as trivial and non-trivial $\mathbb{R}$ -line bundles over $\mathbb{S}^1$, respectively. Both of these bundles come with the Abelian structure group $\text{GL}(1, \mathbb{R}) = \mathbb{R} \setminus \{0\}$. The difference between them is that the cylinder bundle has no negative transition functions $t_{ij} < 0$, resulting in no twists, while the Möbius bundle has one negative transition function, producing a single twist. This raises the question of whether the cylinder and the Möbius bundle are the only two $\mathbb{R}$ -line bundles over $\mathbb{S}^1$. Once more, the answer is provided by the theory of characteristic classes, which confirms that this is the case based on a characteristic class called the first Stiefel–Whitney class $w_1(\mathbb{S}^1)$ of $\mathbb{S}^1$. However, it feels like there is something peculiar going on here. Our intuition tells us that there is no way to continuously deform a Möbius strip with say two twists into a Möbius strip with one twist or a cylinder with no twists. In fact, it seems to be impossible to continuously deform any two Möbius strips with a different number or direction of twists into one another without tearing and glueing. This leads us to expect $\mathbb{Z}$ classes of distinct $\mathbb{R}$ -line bundles, characterised by the number and

direction of twists, instead of the two classes identified by the first Whitney class. Again, the culprit is that our intuition is so attuned to three dimensions: when we attempt to picture distinct $\mathbb{R}$ -line bundles as different configurations of Möbius strips, we are prone to visualising *embeddings* of $\mathbb{R}$ -line bundles in $\mathbb{R}^3$ instead. Mathematically, this corresponds to the distinction between *isomorphism classes* and *isotopy classes* of $\mathbb{R}$ -line bundles. Isomorphism classes are concerned with the number of distinct $\mathbb{R}$ -line bundles as fibre bundles – there is no reference to embedding spaces and there are precisely two such classes characterised by the first Stiefel–Whitney class $w(\mathbb{S}^1)$ of the base space $\mathbb{S}^1$. Isotopy classes, on the other hand, examine the number of distinct embeddings of $\mathbb{R}$ -line bundles into some higher dimensional embedding space A . In $A = \mathbb{R}^3$ there are $\mathbb{Z}$ isotopy classes of $\mathbb{R}$ -line bundles, just like our intuition suggests. The difference between the isotopy classes and isomorphism classes of fibre bundles vanishes when the embedding space becomes sufficiently high-dimensional.¹⁴

Consider a vector bundle with fibre $F = \mathbb{R}^m$. Over every trivialisation chart U , the vector bundle looks like a trivial product, i.e. $\pi^{-1}(U) \simeq U \times \mathbb{R}^m$, so we may choose m linearly independent local sections $\{\sigma_{U,1}, \dots, \sigma_{U,m}\}$ over U . Such a collection of linearly independent local sections is called a local frame. The existence of a global frame is closely related to the topological class of a vector bundle. This is captured by the following theorem.

Theorem 2.2.1. *A vector bundle $E \xrightarrow{\pi} B$ is trivial if and only if it admits a global frame.*

A proof of this theorem can, for instance, be found in Ref. [39].

2.2.2 Principal Bundles

Say we have some base space B and structure group G in mind and we want to construct a simple fibre bundle out of them. Every possible typical fibre has to admit a left action of G so it seems natural to start thinking in terms of representation theory and vector bundles. However, there is a somewhat simpler typical fibre immediately available: the structure group G itself. Recall that G has to be a topological group, so it qualifies as a topological space. Furthermore, G readily comes with a left (and right) action on itself by definition. The result of such a construction is known as a principal G -bundle.

Definition 2.2.4. Principal G -Bundle. A **principal G -bundle** $P \xrightarrow{\pi} B$ is a fibre bundle whose fibre F is identical to the structure group G . A principal G -bundle over B is also often denoted by $P(B, G)$.

By definition, the left (and right) action of G on itself is transitive, i.e. for every $g_1, g_2 \in G$ there exists a group element $l \in G$ (and $r \in G$) such that $g_2 = lg_1$ (and $g_2 = g_1r$). It is common to use the right action of G in the construction of principal bundles, so we will adopt this convention in the following as well. The fibres of a principal G -bundle then correspond to the orbit of the G -action, i.e. for $u \in P(B, G)$ and $\pi(u) = p$ we can construct the fibre $F_p = \pi^{-1}(p)$ as

$$F_p = \{ug \mid g \in G\}. \quad (2.126)$$

This insight will be important for the construction of connections on principal bundles. Given a principal bundle $P(B, G)$ we can use the representation theory of G to generate associated fibre bundles.

Definition 2.2.5. Let $P = P(B, G)$ be a principal G -bundle and let G act on a manifold F on the left. The **associated fibre bundle** $P \times_G F$ is a fibre bundle over the same base space B as P , but with typical fibre F . It is constructed as the quotient $(P \times F)/G$ in which two points $(u, f), (v, h) \in P \times F$ are identified if there exists a $g \in G$ such that

$$(u, f) = g \cdot (v, h), \quad (2.127)$$

where $g \cdot (v, h)$ denotes the induced left action

$$g \cdot (v, h) := (vg, g^{-1}h) \quad (2.128)$$

of G on $P \times F$.

¹⁴This is related to the Whitney embedding theorem which states that every m -dimensional manifold can be embedded in $\mathbb{R}^{2m}$ [69]. An isotopy version of this theorem reveals that for $m \geq 2$ any two embeddings of an m -dimensional manifold into $\mathbb{R}^{2m+1}$ are isotopic. Note that the Whitney embedding theorem and its isotopy version only determine an upper bound for the smallest dimension in which the embedding/isotopy is guaranteed to work. Often, one can do better.

Principal Bundle Theory in Mathematics	Gauge Theory in Physics
Principal bundle	Charge sector
Structure group	Local gauge group
Local trivialisation	Gauge
Choice of local trivialisation	Fixing a gauge
Change of local trivialisation	Local gauge transformation
Local section of associated vector bundle	Matter field
Induced connection on associated vector bundle	Minimal coupling
Connection	Gauge field/potential
Curvature	Gauge field strength

Table 2.2: A translation of concepts between the mathematical theory of principal bundles and physical gauge theory. White rows indicate properties of principal bundles while light gray rows signify properties of associated vector bundles. Table adapted from Ref. [70].

The fibre bundle structure of $E = P \times_G F$ is given as follows. Let $\pi : P \rightarrow B$ denote the projection map of the original principal bundle P . The projection map $\pi_E : E \rightarrow B$ of E is simply defined as $\pi_E(u, f) = \pi(u)$ which is compatible with the equivalence relation because $\pi(u) = \pi(ug)$ for all $g \in G$. The latter is ensured by the fact that the fibres F_p correspond to the orbit of the G -action, cf. Eq. (2.126). The local trivialisations are then given by $\phi_i : U_i \times F \rightarrow \pi^{-1}(U_i)$. Note that the induced left action in Eq. (2.128) is not unique – there exist many other valid left actions on the product $P \times F$. The reason for choosing this one is that it equips the new bundle $P \times_G F$ with the same transition functions of the original bundle $P(B, G)$ [39]. Since the transition functions encode the twisting of a bundle, this ensures that the new bundle $P \times_G F$ is twisted in the same way as the original principal bundle $P(B, G)$. It is in this sense that the new bundle is *associated* to the original one. A particularly useful class of associated fibre bundles are associated vector bundles where we choose a vector space V carrying a representation ρ of G for the new fibre. The concept of associated fibre bundles allows us to translate back and forth between statements about vector bundles and statements about principal bundles.

Principal bundles and associated (vector) bundles form the mathematical basis for gauge theory in physics. The structure group of the principal bundle becomes the local gauge group of the gauge theory; a choice of local trivialisation in the principal bundle amounts to fixing a gauge in the gauge theory. A more detailed translation of concepts between principal bundle theory and physical gauge theory is summarised in Tab. 2.2. While the principal bundles implement the physical gauge fields, the associated vector bundles appear as matter fields in the gauge theory. The bottom rows of Tab. 2.2 show two rather important concepts that we have not yet discussed, namely the connection and curvature of a (principal) bundle representing the gauge potential and the gauge field strength, respectively.

2.2.3 Pullback Bundles and Classifying Spaces

Let $E \xrightarrow{\pi} B$ be a fibre bundle with typical fibre F . Every continuous function $f : B' \rightarrow B$ induces a new fibre bundle over B' with the same typical fibre F .

Definition 2.2.6. Pullback Bundle. Let $E \xrightarrow{\pi} B$ be a fibre bundle with typical fibre F and let $f : B' \rightarrow B$ be a continuous map. The **pullback bundle** f^*E of E by f is defined as

$$f^*E := \{(b', e) \in B' \times E \mid f(b') = \pi(e)\} \subseteq B' \times E \quad (2.129)$$

with the subspace topology and the projection map $\pi' : f^*E \rightarrow B'$ sending $(b', e) \mapsto b'$.

The fiber of f^*E over a point $b' \in B'$ is precisely the fiber of E over $f(b')$, so the pullback bundle copies the fibres of E and attaches them to a new base manifold B' by the continuous map f . Any section $\sigma : B \rightarrow E$ induces a pullback section

$$\begin{aligned} f^*\sigma : B' &\rightarrow f^*E \\ b' &\mapsto (b', (\sigma \circ f)(b')) \end{aligned} \quad (2.130)$$

on f^*E . Similarly, a trivialisation atlas $\{(U_i, \phi_i)\}$ of E induces a trivialisation atlas $\{(V_i, \psi_i)\}$ of f^*E where $V_i = f^{-1}(U_i)$ and where $\psi_i(b', h) = (b', u)$ if $\phi_i(u) = (f(b'), h)$, i.e.

$$\begin{aligned}\psi_i : V_i \times F &\rightarrow \pi'^{-1}(V_i) \\ (b', h) &\mapsto (b', \phi_i^{-1}(f(b'), h)) .\end{aligned}\tag{2.131}$$

Accordingly, the transition functions $t_{ij} : U_i \cap U_j \rightarrow G$ with respect to the trivialisation atlas $\{(U_i, \phi_i)\}$ of E induce the transition functions

$$\begin{aligned}f^*t_{ij} : B' &\rightarrow f^*E \\ b' &\mapsto (t_{ij} \circ f)(b')\end{aligned}\tag{2.132}$$

with respect to the trivialisation atlas $\{(V_i, \psi_i)\}$ on f^*E . It is in this sense that the function $f : B' \rightarrow B$ is used to *pull back* a bundle $E \xrightarrow{\pi} B$ over B to a bundle $f^*E \xrightarrow{\pi'} B'$ over B' .

Pullback bundles play an important role in the classification of principal G -bundles and vector bundles. Specifically, every rank- n vector bundle $E \xrightarrow{\pi} B$ over a compact base space B is isomorphic to a pullback of a universal vector bundle $E_n \xrightarrow{\pi_{E_n}} G_n$ over a so-called classifying space G_n . For real vector bundles, G_n is the real Grassmannian¹⁵

$$G_n \equiv G_n(\mathbb{R}^\infty) ,\tag{2.133}$$

i.e. the smooth manifold of all n -dimensional real linear subspaces of $\mathbb{R}^\infty$. Every point $p \in G_n$ corresponds to an n -dimensional linear subspace $W_p \subset \mathbb{R}^\infty$ so we can define a tautological bundle $E_n \xrightarrow{\pi_{V_n}} G_n$ which attaches to every $p \in G_n$ the particular linear subspace W_p it represents. The tautological bundle $E_n \xrightarrow{\pi_{V_n}} G_n$ is then the universal real vector bundle of all real rank- n vector bundles. Similarly, we can define the Stiefel manifold

$$V_n := V_n(\mathbb{R}^\infty)\tag{2.134}$$

of all orthonormal n -frames in $\mathbb{R}^\infty$. There is a natural projection $V_n \xrightarrow{\pi_{V_n}} G_n$ that sends every n -frame in V_n to the subspace it spans in G_n . The projection $V_n \xrightarrow{\pi_{V_n}} G_n$ defines a principal $O(n)$ -bundle, where the natural right action of $O(n)$ rotates each n -frame of V_n within the subspace it spans. The Stiefel manifold with its natural projection is the universal principal $O(n)$ -bundle. These constructions extend to complex vector bundles and principal $U(n)$ -bundles. Here, the universal complex vector bundle is the tautological bundle of the complex Grassmannian $G_n := G_n(\mathbb{C}^\infty)$ and the universal principal $U(n)$ -bundle is the complex Stiefel manifold $V_n := V_n(\mathbb{C}^\infty)$ with its natural projection to $G_n(\mathbb{C}^\infty)$.

More generally, one can show that every principal G -bundle $P \xrightarrow{\pi_P} B$ is isomorphic to a pullback bundle $f^*EG \xrightarrow{\pi'} B$ of a universal principal G -bundle $EG \xrightarrow{\pi_{EG}} BG$ over a so-called classifying space BG by some continuous function $f : B \rightarrow BG$. The names EG and BG play on the usual notation of E and B for the total and base space, highlighting the distinguished role of *the* universal principal G -bundle EG over *the* classifying space BG . In this terminology, the classifying spaces of principal $O(n)$ - and $U(n)$ -bundles we discussed before are called $BO(n) = G_n(\mathbb{R}^\infty)$ and $BU(n) = G_n(\mathbb{C}^\infty)$, while the total spaces are denoted $EO(n) = V_n(\mathbb{R}^\infty)$ and $EU(n) = V_n(\mathbb{C}^\infty)$, respectively.

The theory of universal bundles and classifying spaces allows us to understand the topology of all principal $O(n)$ - and $U(n)$ -bundles and all real and complex vector bundles in terms of the topology of their respective universal bundles.

¹⁵The real Grassmannian is defined as $G_n(\mathbb{R}^\infty) := \bigcup_{k=n}^\infty G_n(\mathbb{R}^k)$ with the weak limit topology where a set of $G_n(\mathbb{R}^\infty)$ is open iff it intersects every $G_n(\mathbb{R}^k)$ in an open set. The finite-dimensional Grassmannians $G_n(\mathbb{R}^k)$ represent the n -dimensional linear subspaces of $\mathbb{R}^k$ and are often defined as $G_n(\mathbb{R}^k) = O(k)/(O(n) \times O(k-n))$. To understand this construction, take an arbitrary, but fixed n -frame F_n that spans a linear subspace $V_n \subset \mathbb{R}^k$. The set of all n -frames in $\mathbb{R}^k$ is then the orbit of F_n under the orthogonal group $O(k)$ of the ambient space. However, that orbit also contains frames that span the same linear subspace V_n and are in this sense equivalent to F_n . To get rid of these, we have to take the quotient of $O(k)$ by the subgroup $O(V_n)$ that stabilises V_n . That stabiliser subgroup is precisely $O(V_n) = O(n) \times O(k-n)$, i.e. the direct product of the $O(n)$ rotations within the n -dimensional subspace V_n with the $O(k-n)$ rotations in its $(k-n)$ -dimensional complement $V_n^\perp$ within $\mathbb{R}^k$. The complex Grassmannian is defined analogously with $U(n)$ instead of $O(n)$.

2.2.4 Connections on Fibre Bundles

A connection on a fibre bundle $E \xrightarrow{\pi} B$ determines how the local geometric data of the fibres is distributed over the base space B by identifying which points of nearby fibres *correspond* to one another. The need for such a notion becomes evident considering the following question:

Given a vector bundle $\pi : E \xrightarrow{\pi} B$ with typical fibre V , a section $\sigma : B \rightarrow E$ and a vector $X \in T_p B$. What is meant by the directional derivative $d_\sigma X$?

If E is the trivial vector bundle $E = B \times V$ it makes sense to understand the section $\sigma : B \rightarrow B \times V$ as a map $\sigma : B \rightarrow V$ and we can simply write

$$d_\sigma X = \frac{d}{ds} \sigma(\gamma(s))|_{s=0} = \lim_{s \rightarrow 0} \frac{\sigma(\gamma(s)) - \sigma(\gamma(0))}{s} \quad (2.135)$$

for any smooth path $\gamma : [-1, 1] \rightarrow B$ with $\gamma'(0) = X$, thus defining a linear derivative map

$$d_\sigma : T_p B \rightarrow F_p \simeq V. \quad (2.136)$$

However, Eq. (2.135) is only well-defined because $\sigma(\gamma(s))$ and $\sigma(\gamma(0))$ both belong to V and can therefore be added and subtracted. This, in turn, relies on the fact that E is the trivial vector bundle where σ can be viewed as a map $\sigma : B \rightarrow V$. If E is a non-trivial vector bundle, this is no longer the case and Eq. (2.135) ceases to be well-defined. Yet even in a non-trivial vector bundle, the fibres F_p and F_q over different points $p \neq q \in B$ are isomorphic so the definition of d_σ as a linear map $T_p B \rightarrow F_p$ should still be possible. The problem is that there is no natural isomorphism between F_p and F_q for $p \neq q$ in non-trivial bundles. We are missing a piece of extra structure that connects these fibres in a unique way, at least if p and q are sufficiently close. That missing piece of extra structure is called a *connection*.

There are various ways to define a connection, but one in particular stands out for its conceptual generality: an *Ehresmann connection* establishes the desired relation between neighbouring fibres by purely geometrical means. Let $E \xrightarrow{\pi} B$ be a smooth fibre bundle. The basic idea of an Ehresmann connection is to identify so-called *vertical* and *horizontal* directions in E that are pointing *along* and *across* fibres, respectively. At any point $u \in E$, moving in the vertical direction means to stay within the current fibre while moving in the horizontal direction means to smoothly change between neighbouring fibres. The latter allows us to identify points between neighbouring fibres as those elements that are horizontally connected. In a local trivialisation around a point $p \in U \subset B$ we may write any $u \in \pi^{-1}(U) \subset E$ as $u \simeq (p, f) \in U \times F$ and identify $f \in F_p$ with all $h \in F_q$ that can be reached by spreading out from (p, f) in the horizontal direction. The notion of a horizontal layer of equivalent points across fibres has a rather similar flavour to the idea of a section. So what is the difference? A section is a “static” slice through the bundle, providing a fixed choice of points in each fibre. In contrast, the horizontal identification of points provided by an Ehresmann connection is a “dynamic” process that involves the transport of elements between distinct fibers. This “dynamic” nature of the connection is captured by defining the Ehresmann connection in terms of vectorial objects – namely elements of the tangent bundle of the original fibre bundle. In the following, we prepare the formal definition of the Ehresmann connection.

In differential geometry, any smooth map $f : M \rightarrow N$ between smooth manifolds induces two mutually dual maps, the push-forward $f_* : TM \rightarrow TN$ and the pullback $f^* : T^*N \rightarrow T^*M$ by f , where the names signify whether the respective maps point in the same (push-forward) or in the opposite (pullback) direction of f . Note that the pullback of differential geometry is related, but not equivalent to the pullback of fibre bundles we discussed before. Let M, N be smooth manifolds and let $f : M \rightarrow N$ be a smooth map. The push-forward f_* of f is a map

$$f_* : TM \rightarrow TN, \quad (2.137)$$

sending elements of the tangent bundle $TM \xrightarrow{\pi_{TM}} M$ of M to elements of the tangent bundle $TN \xrightarrow{\pi_{TN}} N$ of N . It is defined pointwise as

$$\begin{aligned} f_* : T_p M &\rightarrow T_{f(p)} N \\ X &\mapsto f_*(X). \end{aligned} \quad (2.138)$$

The elements $X \in T_p M$ are tangent vectors of M at p . They can be defined as equivalence classes $[\gamma]$ of smooth curves $\gamma : [-1, 1] \rightarrow M$ with $\gamma(0) = p$ and coinciding derivatives in M . In that case, we write $X = [\gamma] \equiv \gamma'(0)$ and $f_*(X) \in T_{f(p)} N$ is given by

$$f_*(\gamma'(0)) := (f \circ \gamma)'(0), \quad (2.139)$$

i.e. the push-forward of a tangent vector $X \equiv \gamma'(0)$ to the curve $\gamma(s)$ at $s = 0$ in M is a tangent vector $f_*(X) \equiv (f \circ \gamma)'(0)$ to the curve $(f \circ \gamma)(s)$ at $s = 0$ in N . The pullback f^* of f is a map

$$f^* : T^* N \rightarrow T^* M, \quad (2.140)$$

sending elements of the cotangent bundle $T^* N \xrightarrow{\pi_{T^* N}} N$ of N to elements of the cotangent bundle $T^* M \xrightarrow{\pi_{T^* M}} M$ of M . It is defined pointwise as

$$\begin{aligned} f^* : T_{f(p)}^* N &\rightarrow T_p^* M \\ \omega &\mapsto f^*(\omega). \end{aligned} \quad (2.141)$$

The elements $\omega \in T_{f(p)}^* N$ are linear functions $\omega : T_{f(p)} N \rightarrow \mathbb{R}$. Therefore, it is natural to take a test vector $X \in T_p M$ and define the pullback $f^*(\omega) \in T_p^* M$ of $\omega \in T_{f(p)}^* N$ via its dual, the push-forward, as

$$(f^*(\omega))(X) = \omega(f_*(X)), \quad (2.142)$$

or, more compactly, as

$$\langle f^*\omega, X \rangle = \langle \omega, f_*X \rangle, \quad (2.143)$$

using a bracket notation $\omega(X) = \langle \omega, X \rangle$ for the dual pairing of linear forms and vectors and omitting the argument brackets for better readability.

A formal definition of the Ehresmann connection can then be given in terms of the push-forward π_* of the projection map. Let $E \xrightarrow{\pi} B$ be a smooth fibre bundle and let $\pi_* : TE \rightarrow TB$ be the push-forward of the projection map π of E . The kernel

$$VE := \ker(\pi_* : TE \rightarrow TB) \quad (2.144)$$

of π_* defines a fibre bundle called the vertical bundle $VE \xrightarrow{\pi_{VE}} E$ over E . The vertical bundle VE is canonically defined for every fibre bundle E . It forms a smooth subbundle of the tangent bundle TE and consists of vectors $X \in TE$ that are *tangent to the fibres* in that they are collapsed by π_* .

Definition 2.2.7. Ehresmann Connection. Let $E \xrightarrow{\pi} B$ be a smooth fibre bundle. An **Ehresmann connection** on E is a smooth subbundle $HE \xrightarrow{\pi_{HE}} E$ of the tangent bundle $TE \xrightarrow{\pi_{TE}} E$, called the **horizontal bundle** of the connection, which is *complementary* to the canonical vertical bundle $VE \xrightarrow{\pi_{VE}} E$ in the sense that

$$TE = VE \oplus HE. \quad (2.145)$$

The projection maps π_{HE} and π_{VE} are the restrictions of the canonical projection map π_{TE} of the tangent bundle.

Equation (2.145) requires that the tangent bundle TE splits as the direct sum of its smooth subbundles VE and HE . For a smooth fibre bundle $E \xrightarrow{\pi} B$ with fibre F this essentially means that at each point $u \in E$ the tangent space $T_u E$ can be written as the direct sum

$$T_u E = V_u E \oplus H_u E \quad (2.146)$$

of a $\dim_{\mathbb{R}}(F)$ -dimensional linear subspace $V_u E \subset T_u E$ called the vertical subspace and a $\dim_{\mathbb{R}}(B)$ -dimensional linear subspace $H_u E \subset T_u E$ called the horizontal subspace. While Eq. (2.144) shows that the vertical subspace $V_u E$ is uniquely defined by the projection map, there is a considerable degree of freedom in the choice of the horizontal subspace $H_u E$ and, by extension, the Ehresmann connection.

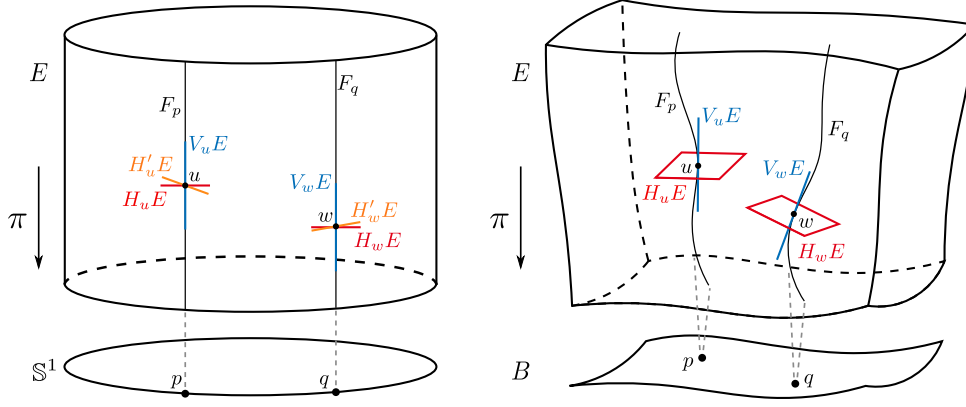


Figure 2.8: Possible choices of horizontal subspaces complementing the vertical subspaces at different points $u \in \pi^{-1}(p) = F_p$ and $w \in \pi^{-1}(q) = F_q$ in two simple fibre bundles. The left picture shows horizontal subspaces of two different Ehresmann connections HE (red) and $H'E$ (orange) on the two-dimensional cylinder bundle $E \xrightarrow{\pi} \mathbb{S}^1$, while the right picture illustrates the concept for a more generic three-dimensional bundle $E \xrightarrow{\pi} B$ with two-dimensional base manifold B and one-dimensional typical fibre F . Illustration created by the author, based on Ref. [71].

This is illustrated in the left picture of Fig. 2.8 for the trivial cylinder bundle $E = \mathbb{S}^1 \times \mathbb{R}$. Embedded in $\mathbb{R}^3$ we can write every point $u \in E$ as $u = (z, e^{i\varphi}) \equiv (z, \varphi)$ using standard cylinder coordinates. Elements of the tangent space $T_u E$ at any $u \in E$ are then vectors

$$X = X_z \frac{\partial}{\partial z} + X_\varphi \frac{\partial}{\partial \varphi} \quad (2.147)$$

with coefficients $X_z, X_\varphi \in \mathbb{R}$. The vertical subspace $V_u E$ at any $u \in E$ is then given by $V_u E = \text{span}(\partial/\partial z)$. For the horizontal subspace we may choose the canonic option $H_u E := \text{span}(\partial/\partial \varphi)$, shown as red lines in the left picture of Fig. 2.8. Clearly, an element of the direct sum $V_u E \oplus H_u E$ is of the form Eq. (2.147) which makes the canonic option $H_u E$ a valid choice for a horizontal subspace. However, we may also choose $H'_u E := \text{span}(\partial/\partial \varphi + f(z, \varphi)\partial/\partial z)$ for any smooth function $f : E \rightarrow \mathbb{R}, (z, \varphi) \mapsto f(z, \varphi)$. Elements of the direct sum $V_u E \oplus H'_u E$ are of the form

$$\begin{aligned} X &= X_V e_V + X_{H'} e_{H'} \\ &= X_V \frac{\partial}{\partial z} + X_{H'} \left(\frac{\partial}{\partial \varphi} + f(z, \varphi) \frac{\partial}{\partial z} \right) \\ &= (X_V + X_{H'} f(z, \varphi)) \frac{\partial}{\partial z} + X_{H'} \frac{\partial}{\partial \varphi}, \end{aligned} \quad (2.148)$$

where $X_V, X_{H'} \in \mathbb{R}$ denote the coefficients in the basis vectors e_V of $V_u E$ and $e_{H'}$ of $H'_u E$. Since the coefficients are arbitrary real numbers, this is equivalent to Eq. (2.147) as well. The horizontal spaces of such an Ehresmann connection are illustrated as orange lines in the left picture of Fig. 2.8. Since the tangent bundle $TE \xrightarrow{\pi_{TE}} E$ of a given bundle $E \xrightarrow{\pi} B$ is a smooth manifold of dimension $\dim_{\mathbb{R}}(TE) = 2\dim_{\mathbb{R}}(E)$, the degree of freedom in the definition of an Ehresmann connection becomes large quickly. This can be recognised from the right picture of Fig. 2.8, which visualises a local chunk of a three-dimensional bundle $E \xrightarrow{\pi} B$ with two-dimensional base manifold B and one-dimensional typical fibre F . The degree of freedom in defining the horizontal subspaces $H_u E$ and $H_w E$ is even larger than in the cylinder example.

Another, equivalent way to define an Ehresmann connection is by using a projection map Φ onto the vertical bundle VE . Such a projection map is a vector bundle endomorphism $\Phi : TE \rightarrow TE$ satisfying the two properties $\Phi^2 = \Phi$ and $\Phi|_{VE} = \text{id}_{VE}$. The former makes Φ a projection, while the latter ensures

$$\text{im}(\Phi) = VE \quad \text{and} \quad \ker(\Phi) = HE, \quad (2.149)$$

which reproduces the Ehresmann separation $TE = \text{im}(\Phi) \oplus \ker(\Phi) = VE \oplus HE$ of TE into vertical and horizontal subbundles. At each point $u \in E$, the projection map $\Phi : T_u E \rightarrow T_u E$ is a linear mapping from

$T_u E$ to itself. In this sense, Φ acts like a $T_u E$ -valued linear form at every $u \in E$. Since Φ defines a smooth distribution of such $T_u E$ -valued linear forms over E , it may be regarded as a TE -valued one-form on E . For this reason, Φ is often called the connection form of the Ehresmann connection. This perspective on the (Ehresmann) connection is very useful in many situations and often present in the description of topological phenomena in physics.

There are various types of connections for different kinds of fibre bundles. Examples include linear (or Koszul) connections on vector bundles, metric connections, like the Riemannian and the Levi-Civita connection, on metric vector bundles, and principal connections on principal bundles. Although these different connection types are generally defined in different ways, they can often be regarded as special cases of certain Ehresmann connections. Due to its great importance for physics, we will take a closer look at one such special case in the following: the principal connection of a principal bundle.

Definition 2.2.8. Principal Connection. Let $E = P(B, G)$ be a principal G -bundle. An Ehresmann connection HE on E is called a **principal connection** if it is invariant under the G -action on E .

This means the following. Remember that a principal G -bundle E admits a right action of G which is formally a map

$$\begin{aligned} R_g : E &\rightarrow E \\ u &\mapsto R_g(u) := ug, \end{aligned} \tag{2.150}$$

where $g \in G$ is arbitrary and fixed. The right action induces a push-forward

$$\begin{aligned} R_{g*} : T_u E &\rightarrow T_{\lambda u} E \\ X &\mapsto R_{g*}(X) \end{aligned} \tag{2.151}$$

between the tangent spaces at every point $u \in E$. In Eq. (2.126) we identified the fibres of E as the orbits of the G -action. Therefore, $u \in E$ and $ug \in E$ correspond to the same fibre, i.e. $\pi(u) = \pi(ug)$, for every $g \in G$. The Ehresmann connection HE of E is said to be invariant under the G -action on E if the horizontal subspaces $H_u E$ and $H_{ug} E$ on the same fibre are related by

$$H_{ug} E = R_{g*} H_u E \tag{2.152}$$

for all $g \in G$ and $u \in E$. As a consequence, a horizontal subspace $H_u E$ at any $u \in E$ generates all the horizontal subspaces on the same fibre. If G is a Lie group, which it usually is in physical applications, we can use the fact that G acts vertically on E to show that the connection form Φ of the Ehresmann connection can be viewed as a one-form ω on E with values in the Lie algebra $\mathfrak{g}$ of G .¹⁶ Formally, we write $\omega \in \Gamma(T^* E \otimes \mathfrak{g}, E)$, which means that ω is a smooth section of the bundle $T^* E \otimes \mathfrak{g}$ over E , i.e. of the cotangent bundle $T^* E$ over E with values in $\mathfrak{g}$. The latter is signified by the tensor product between $T^* E$ and $\mathfrak{g}$. Based on ω we can define an important object called the local connection form $\mathcal{A}_i$ of the connection. The important difference between $\mathcal{A}_i$ and ω is that $\mathcal{A}_i$ is a $\mathfrak{g}$ -valued one-form on an open neighbourhood $U_i \subset B$ of the *base manifold* B of E , rather than one on the *entire bundle* E itself. In practice, this makes $\mathcal{A}_i$ much more accessible than ω or Φ . In fact, the local gauge potential from gauge theory corresponds to neither Φ nor ω , but to a local connection form $\mathcal{A}_i$. A general definition of $\mathcal{A}_i$ can be given in terms of the pullback $\sigma_i^* : T^* E \rightarrow T^* U_i$ of a local section $\sigma_i : U_i \rightarrow E$.

Definition 2.2.9. Local Connection Form of a Principal Bundle. Let $E = P(B, G)$ be a principal G -bundle with an Ehresmann connection HE and let $\omega \in \Gamma(T^* E \otimes \mathfrak{g}, E)$ be the connection one-form of HE . The **local connection form** $\mathcal{A}_i$ is defined as the pullback

$$\mathcal{A}_i = \sigma_i^*(\omega) \tag{2.153}$$

of ω by a local section $\sigma_i : U_i \rightarrow E$ of E , making it an element $\mathcal{A}_i \in \Gamma(T^* U_i \otimes \mathfrak{g}, U_i)$.

Note that the tensor product with $\mathfrak{g}$ does not affect the construction. Strikingly, a local connection form $\mathcal{A}_i$ on a trivialisation chart $U_i \subset B$ completely determines the connection form ω_i on $\pi^{-1}(U_i) \subset E \xrightarrow{\pi} B$,

¹⁶This is done using the push-forward $R_{\varphi*}$ of the action R_{φ} of the one-parameter subgroups $\varphi : \mathbb{R} \rightarrow G$ of G [39].

although it *does* depend on the particular local section $\sigma_i : U_i \rightarrow E$. Consider an open cover $\{U_i\}$ of the base space B of our principal bundle E together with local sections $\{\sigma_i\}$ and local connection forms $\{\mathcal{A}_i\}$ defined on the open neighbourhoods of that cover. The local sections $\{\sigma_i\}$ and local connection forms $\{\mathcal{A}_i\}$ on the $\{U_i\}$ patches uniquely determine local connection one-forms $\{\omega_i\}$ on the chunks $\{\pi^{-1}(U_i)\}$ of E located above these patches. For ω to be uniquely defined on all of E , the local forms ω_i and ω_j must agree on all nonempty overlaps $U_i \cap U_j \neq \emptyset$, ensuring that the collection $\{\omega_i\}$ pieces together into a single well-defined global form ω with $\omega|_{U_i} = \omega_i$. To fulfil this condition, the local connection forms $\mathcal{A}_i$ have to satisfy the transformation property [39]

$$\mathcal{A}_j = t_{ij}^{-1} \mathcal{A}_i t_{ij} + t_{ij}^{-1} dt_{ij} , \quad (2.154)$$

where $t_{ij} : U_i \cap U_j \rightarrow G$ denote the transition functions between the local sections σ_i over U_i and σ_j over U_j . The first term on the right-hand side of Eq. (2.154) corresponds to the standard adjoint action $\text{Ad}_g : \mathfrak{g} \rightarrow \mathfrak{g}, \mu \mapsto g^{-1} \mu g$ of the Lie group G on its Lie algebra $\mathfrak{g}$. The second term is best understood remembering that at each point $p \in U_i \cap U_j$ the local connection forms $\mathcal{A}_i$ are maps $\mathcal{A}_i : T_p(U_i \cap U_j) \rightarrow \mathfrak{g}$ taking a vector $X \in T_p(U_i \cap U_j)$ to an element $\mu \in \mathfrak{g}$ of the Lie algebra. Acting on a test vector $X \in T_p(U_i \cap U_j)$ at $p \in U_i \cap U_j$ it becomes

$$(t_{ij}^{-1}(p) dt_{ij})(X) = t_{ij}^{-1}(p) \frac{d}{ds} t_{ij}(\gamma(s)) \Big|_{s=0} = \frac{d}{ds} [t_{ij}^{-1}(p) t_{ij}(\gamma(s))] \Big|_{s=0} , \quad (2.155)$$

where $\gamma : [-1, 1] \rightarrow U_i \cap U_j$ is a curve with $\gamma(0) = p$ and $\gamma'(0) = X$. Note that $\gamma(0) = p$ ensures that Eq. (2.155) is an element of the tangent space $T_e G \simeq \mathfrak{g}$ of G at the identity e since it corresponds to a derivative at the identity $t_{ij}^{-1}(p) t_{ij}(\gamma(0)) = t_{ij}^{-1}(p) t_{ij}(p) = e$ by definition.

As a consequence, we can construct ω over E from any open cover $\{U_i\}$ of B with local sections $\{\sigma_i\}$ and local connection forms $\{\mathcal{A}_i\}$ given the $\mathcal{A}_i$ satisfy Eq. (2.154). However, there is an important caveat here, namely that non-trivial principal bundles do not admit a global section. So while $\mathcal{A}_i = \sigma_i^* \omega$ always exists locally, it may not be defined globally. These observations have immediate implications for physics. We mentioned earlier that the local connection form $\mathcal{A}_i$ corresponds to the gauge potential in physical gauge theories. Consequently, the compatibility condition in Eq. (2.154) becomes a gauge transformation in gauge theory (cf. Tab. 2.2) and the non-triviality of a principal bundle manifests in the form of topological gauge field configurations, such as instantons or monopoles.

An Ehresmann connection on a smooth fibre bundle E also determines a generalised notion of parallel transport known as the horizontal lift. The horizontal lift specifies how a curve $\gamma : [0, 1] \rightarrow B$ in the base space B of a fibre bundle E can be lifted to a curve $\tilde{\gamma} : [0, 1] \rightarrow E$ in E .

Definition 2.2.10. Horizontal Lift. Let $E \xrightarrow{\pi} B$ be a smooth fibre bundle with an Ehresmann connection $HE \xrightarrow{\pi_{HE}} E$ and let $\gamma : [0, 1] \rightarrow B$ be a smooth curve in B . A **horizontal lift** $\tilde{\gamma}$ of γ is a curve $\tilde{\gamma} : [0, 1] \rightarrow E$ satisfying $(\pi \circ \tilde{\gamma}) = \gamma$ and $\tilde{\gamma}'(s) \in H_{\tilde{\gamma}(s)} E$ for all $s \in [0, 1]$.

The first condition $(\pi \circ \tilde{\gamma}) = \gamma$ means that $\tilde{\gamma}$ is projected onto the original curve γ so it is indeed a version of γ that has only been “lifted upwards”. The second condition asks that $\tilde{\gamma}'(s) \in H_{\tilde{\gamma}(s)} E$ everywhere so the tangent vectors to $\tilde{\gamma}$ are required to be always horizontal. In terms of the connection form Φ , the requirement $\tilde{\gamma}'(s) \in H_{\tilde{\gamma}(s)} E$ becomes

$$\Phi(\tilde{\gamma}'(s)) = 0 . \quad (2.156)$$

Since Φ can be understood as a differential one-form, Eq. (2.156) constitutes an ordinary differential equation (ODE) and the fundamental theorem of ODEs guarantees the local existence and uniqueness of the horizontal lift [39].

Theorem 2.2.2. Uniqueness of Horizontal Lifts. Let $\gamma : [0, 1] \rightarrow B$ be a smooth curve in the base space B of a smooth fibre bundle $E \xrightarrow{\pi} B$ with an Ehresmann connection and let $u_0 \in \pi^{-1}(\gamma(0))$ be an arbitrary, but fixed element of the fibre over $\gamma(0)$. Then there exists a unique horizontal lift $\tilde{\gamma}(s)$ in E such that $\tilde{\gamma}(0) = u_0$.

The requirement that all tangent vectors must be horizontal means that $\tilde{\gamma}$ follows a path in E that only connects elements of neighbouring fibres that are considered equivalent by the Ehresmann connection.

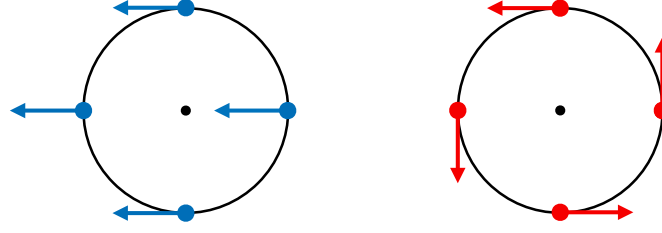


Figure 2.9: Parallel transport of tangent vectors to the punctured plane $\mathbb{R}_*^2 := \mathbb{R}^2 \setminus \{0\}$ along the closed path $\mathbb{S}^1 \subset \mathbb{R}_*^2$ according to two different Levi-Civita connections on the tangent bundle $T\mathbb{R}_*^2$. The parallel transport in the left picture is defined using the Levi-Civita connection of the standard metric $ds^2 = dx^2 + dy^2 = dr^2 + r^2 d\theta^2$, while the right picture is defined using the Levi-Civita connection of the modified metric $ds^2 = dr^2 + d\theta^2$.

In this sense, the horizontal lift $\tilde{\gamma}$ describes a transport of $\tilde{\gamma}(0)$ without change, i.e. a parallel transport of $\tilde{\gamma}(0)$ similar to the familiar parallel transport of vectors in tangent bundles. For this reason we call $\tilde{\gamma}(1) \in \pi^{-1}(\gamma(1))$ the parallel transport of $\tilde{\gamma}(0) \in \pi^{-1}(\gamma(0))$ along γ . The uniqueness of horizontal lifts ensures that for any given curve $\gamma : [0, 1] \rightarrow B$ and $u_0 \in \pi^{-1}(\gamma(0))$ there exists a unique parallel transport $u_1 \in \pi^{-1}(\gamma(1))$ of u_0 along γ that we obtain by lifting up $\gamma(0)$ to $\tilde{\gamma}(0) := u_0$ and setting $u_1 := \tilde{\gamma}(1)$. This naturally defines a map

$$P_\gamma : \pi^{-1}(\gamma(0)) \rightarrow \pi^{-1}(\gamma(1))$$

$$u_0 \mapsto u_1, \quad (2.157)$$

called the parallel transport P_γ along γ . Note that P_γ does not only depend on the path γ , but also on the specific connection chosen on the bundle. This is illustrated in Fig. 2.9, which shows the parallel transport of tangent vectors to the punctured plane $\mathbb{R}_*^2 := \mathbb{R}^2 \setminus \{0\}$ along the unit circle $\mathbb{S}^1 \subset \mathbb{R}_*^2$ for two different Levi-Civita connections on the tangent bundle $T\mathbb{R}_*^2$. Levi-Civita connections are a special class of metric connections that are uniquely defined for any given (pseudo-)Riemannian metric on a smooth Riemannian manifold. The two metrics that give rise to the two distinct Levi-Civita connections on $T\mathbb{R}_*^2$ illustrated in Fig. 2.9 are the standard metric $ds^2 = dx^2 + dy^2 = dr^2 + r^2 d\theta^2$ on the left, and the modified metric $ds^2 = dr^2 + d\theta^2$ on the right. The latter is singular at the origin, which is why we consider $\mathbb{R}_*^2$ instead of $\mathbb{R}^2$. Even though Levi-Civita connections are rather important for physics, we are not going to go into more detail about them.

In a principal G -bundle $E = P(B, G)$, we can use a version of the compatibility condition Eq. (2.154) to determine the parallel transport u_1 explicitly. Recall that Eq. (2.154) describes the transformation behaviour of local connection forms $\mathcal{A}_i$ and $\mathcal{A}_j$ on overlapping neighbourhoods $U_i \cap U_j \neq \emptyset$ with local sections σ_i and σ_j . Since E is a principal bundle, σ_i and σ_j are related as

$$\sigma_j(p) = \sigma_i(p)t_{ij}(p) \quad (2.158)$$

with $t_{ij}(p) \in G$ for all $p \in U_i \cap U_j$. In fact, this is why the transition functions t_{ij} appear in Eq. (2.154) to begin with. If we take a chart $U_i \subset B$ that supports a given curve $\gamma : [0, 1] \rightarrow U_i \subset B$ and a local section $\sigma_i : U_i \rightarrow \pi^{-1}(U_i) \subset E$, we may write the horizontal lift $\tilde{\gamma} : [0, 1] \rightarrow \pi^{-1}(U_i) \subset E$ in a similar way, namely as

$$\tilde{\gamma}(s) = \sigma_i(\gamma(s))g_i(\gamma(s)), \quad (2.159)$$

where $g_i(\gamma(s)) \in G$ everywhere. Therefore, $\tilde{\gamma}$ behaves just like a section σ_j and we can understand the compatibility condition Eq. (2.154) as a relation between $\mathcal{A}_j = \sigma_j^*(\omega) = \tilde{\gamma}^*(\omega)$ and $\mathcal{A}_i = \sigma_i^*(\omega)$. If we plug $\mathcal{A}_j = \tilde{\gamma}^*(\omega)$ into Eq. (2.154) and apply it to $X = \gamma'(s)$ we get

$$(\tilde{\gamma}^*(\omega))(X) = g_i^{-1}(s)\mathcal{A}_i(X)g_i(s) + (g_i^{-1}(s)dg_i)(X), \quad (2.160)$$

where we wrote $g_i^{(-1)}(s) \equiv g_i^{(-1)}(\gamma(s))$. The dual relationship between the pullback and the push-forward as defined in Eq. (2.142) allows us to rephrase this as

$$\omega(\tilde{\gamma}_*(X)) = g_i^{-1}(s)\mathcal{A}_i(X)g_i(s) + (g_i^{-1}(s)dg_i)(X). \quad (2.161)$$

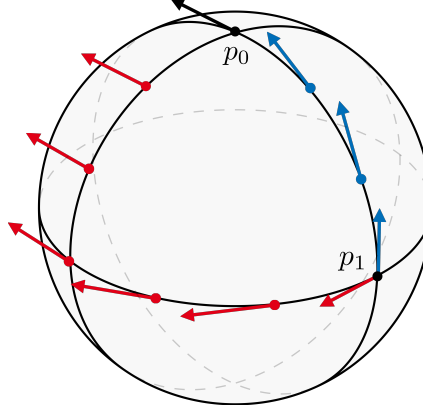


Figure 2.10: Parallel transport of tangent vectors on $\mathbb{S}^2$. The black tangent vector at the north pole p_0 is parallel transported to p_1 at the equator along two different paths indicated by red and blue arrows, respectively. The mismatch between the red and blue tangent vectors at p_1 is a holonomy of the tangent bundle $T\mathbb{S}^2$ of $\mathbb{S}^2$ and famously attributed to the curvature of $\mathbb{S}^2$.

Now, the push-forward $\tilde{X} = \tilde{\gamma}_*(X)$ of the tangent vector $X = \gamma'(0)$ to γ by the horizontal lift $\tilde{\gamma}$ is itself a tangent vector to $\tilde{\gamma}$. By definition of the horizontal lift, $\tilde{X}$ is horizontal and hence collapsed as $\omega(\tilde{X}) = 0$ by the connection form ω on E , cf. Eq. (2.156). With this, Eq. (2.161) becomes

$$0 = g_i^{-1}(s)\mathcal{A}_i(X)g_i(s) + (g_i^{-1}(s)dg_i)(X). \quad (2.162)$$

If we rewrite the second term on the right-hand side according to the middle expression of Eq. (2.155) we end up with

$$0 = g_i^{-1}(s)\mathcal{A}_i(X)g_i(s) + g_i^{-1}(s)\frac{d}{ds}g_i(\gamma(s))\Big|_{s=0}, \quad (2.163)$$

which immediately rearranges into

$$\frac{d}{ds}g_i(\gamma(s)) = -\mathcal{A}_i(X)g_i(s) \quad (2.164)$$

after multiplying by $g_i(s)$ from the left. For $g_i(0) := e$ the formal solution to Eq. (2.164) reads

$$g_i(\gamma(s)) = \mathcal{P} \exp \left[- \int_{\gamma(0)}^{\gamma(1)} \mathcal{A}_{i,\mu}(\gamma(s)) dx^\mu \right], \quad (2.165)$$

where $\mathcal{P}$ is the path ordering operator along $\gamma(s)$ and where $\mathcal{A}_{i,\mu}$ are the components of $\mathcal{A}_i$ in local coordinates x^μ on $U_i \subset B$. Using this expression for $g_i(\gamma(s))$ and the definition Eq. (2.159) of the horizontal lift, we find that the parallel transport u_1 takes the form

$$u_1 = \sigma_i(\gamma(1))\mathcal{P} \exp \left[- \int_{\gamma(0)}^{\gamma(1)} \mathcal{A}_{i,\mu}(\gamma(s)) dx^\mu \right]. \quad (2.166)$$

Versions of this formula are ubiquitous in the physics literature. Most notably, it appears in the form of an observable called the Berry phase, which is a special case of an important notion called holonomy.

2.2.5 Holonomy

Let $E \xrightarrow{\pi} B$ be a smooth fibre bundle with Ehresmann connection HE . Say we take two distinct curves $\alpha : [0, 1] \rightarrow B$ and $\beta : [0, 1] \rightarrow B$ with $\alpha(0) = \beta(0) = p_0$ and $\alpha(1) = \beta(1) = p_1$ and determine their horizontal lifts $\tilde{\alpha} : [0, 1] \rightarrow E$ and $\tilde{\beta} : [0, 1] \rightarrow E$ such that $\tilde{\alpha}(0) = \tilde{\beta}(0) = u_0 \in \pi^{-1}(p_0)$. Then $\tilde{\alpha}(1)$ is not always equal to $\tilde{\beta}(1)$. Figure 2.10 shows a prominent example of this: the parallel transport of tangent vectors on $\mathbb{S}^2$ leads to very different results depending on the specific path of transport. Generally, if we have two curves α and β with $\alpha(0) = \beta(0) = p_0$ and $\alpha(1) = \beta(1) = p_1$ we may concatenate them as

$$\gamma(s) = \begin{cases} \alpha(1-2s) & \text{for } s \in [0, \frac{1}{2}] \\ \beta(2s-1) & \text{for } s \in [\frac{1}{2}, 1] \end{cases} \quad (2.167)$$

to get a loop in B . Since we chose $\tilde{\alpha}(0) = \tilde{\beta}(0) = u_0$, the horizontal lift $\tilde{\gamma}$ of γ can be written as

$$\tilde{\gamma}(s) = \begin{cases} \tilde{\alpha}(1 - 2s) & \text{for } s \in [0, \frac{1}{2}) \\ \tilde{\beta}(2s - 1) & \text{for } s \in [\frac{1}{2}, 1] \end{cases}, \quad (2.168)$$

which connects $\tilde{\gamma}(0) = \tilde{\alpha}(1) =: u_\alpha$ and $\tilde{\gamma}(1) = \tilde{\beta}(1) =: u_\beta$. The start point u_α and end point u_β of $\tilde{\gamma}$ are still in the same fibre as $\pi(u_\alpha) = \pi(u_\beta) = p_1$, but whenever $u_\alpha \neq u_\beta$ the lift $\tilde{\gamma}$ is no longer a loop. In this sense, every loop γ in the base space B defines a non-trivial map

$$\begin{aligned} \tau_\gamma : \pi^{-1}(p) &\rightarrow \pi^{-1}(p) \\ u &\mapsto \tilde{\gamma}_u, \end{aligned} \quad (2.169)$$

that takes an element $u \in F_p$ to another element $\tilde{\gamma}_u \in F_p$ defined as the end point $\tilde{\gamma}_u := \tilde{\gamma}_u(1)$ of the unique horizontal lift $\tilde{\gamma}_u$ that starts at $\tilde{\gamma}_u(0) = u$. The τ_γ is completely determined by the loop γ in B and the connection HE on E . The fact that $u, \tau_\gamma(u) \in F_p$ means that they are related as

$$\tau_\gamma(u) = ug_\gamma(u) \quad (2.170)$$

for some $g_\gamma(u) \in G$. This motivates the following definition.

Definition 2.2.11. Holonomy Group. Let $E \xrightarrow{\pi} B$ be a smooth fibre bundle with structure group G and Ehresmann connection form Φ . Take an element $u \in E$ with $\pi(u) = p$ and consider the set $C_p(B) = \{\gamma : [0, 1] \rightarrow B \mid \gamma(0) = \gamma(1) = p\}$ of loops in B based at p . The set

$$\text{Hol}_u(\Phi) = \{g \in G \mid \exists \gamma \in C_p(B) : \tau_\gamma(u) = ug\} \quad (2.171)$$

defines a subgroup of the structure group G that is called the **holonomy group** of Φ at $u \in E$.

Note that the order in which we concatenate α and β in Eq. (2.167) only changes the orientation of the resulting loop, but not its end point. In contrast, the order in which we concatenate the horizontal lifts $\tilde{\alpha}$ and $\tilde{\beta}$ in Eq. (2.168) does have an impact on the respective end point. The holonomy group therefore measures the failure of parallel transport to *commute*. This indicates that the holonomy group generally contains sensitive geometrical information – its elements depend strongly on the connection of the fibre bundle and the geometric details of the underlying loops in the base manifold. However, there is also some topological information to be found in the holonomy group. To see this, recall that we used homotopy-equivalence classes of the same kind of one-dimensional loops to define the fundamental group π_1 of a topological space earlier. For example, we found that the identity of the fundamental group corresponds to contractible loops, i.e. loops that can be shrunk down to a point representing the trivial constant loop there. We can translate this idea to the holonomy group. In fact, we find that the subset $\text{Hol}_u^0(\Phi)$ of the holonomy group that comes from contractible loops in B forms a connected normal subgroup $\text{Hol}_u^0(\Phi) \triangleleft \text{Hol}_u(\Phi)$ that is called the restricted holonomy group of Φ at $u \in E$. As such it contains the identity element of $\text{Hol}_u(\Phi)$, namely the constant loop $\gamma_0(s) = \pi(u)$, which makes $\text{Hol}_u^0(\Phi)$ the identity component of $\text{Hol}_u(\Phi)$. Since $\text{Hol}_u^0(\Phi)$ is a normal subgroup of $\text{Hol}_u(\Phi)$, the quotient

$$\text{Hol}_u(\Phi)/\text{Hol}_u^0(\Phi) = \{g\text{Hol}_u^0(\Phi) \mid g \in \text{Hol}_u(\Phi)\} \quad (2.172)$$

has a group structure describing a partition of $\text{Hol}_u(\Phi)$ with respect to $\text{Hol}_u^0(\Phi)$. A more comprehensive understanding of the topological information contained in $\text{Hol}_u(\Phi)$ is then provided by a natural surjective group homomorphism

$$H : \pi_1(B) \rightarrow \text{Hol}_u(\Phi)/\text{Hol}_u^0(\Phi) \quad (2.173)$$

between the fundamental group $\pi_1(B)$ of the base manifold and the quotient $\text{Hol}_u(\Phi)/\text{Hol}_u^0(\Phi)$ of the bundle holonomy group by its restricted holonomy group. The fact that H is surjective tells us that it takes at least one element of $\pi_1(B)$ to each element of $\text{Hol}_u(\Phi)/\text{Hol}_u^0(\Phi)$. This allows us to identify two extreme cases. The first one arises when $\pi_1(B) = \{e\}$, i.e. when the fundamental group of the base manifold B is trivial so B is simply connected. In that case, the surjectivity of H tells us that

$\text{Hol}_u(\Phi)/\text{Hol}_u^0(\Phi) \simeq \{e\}$, which immediately implies that $\text{Hol}_u(\Phi) = \text{Hol}_u^0(\Phi)$. The holonomy group over simply connected manifolds is therefore guaranteed to be connected (only has the connected component of the identity). According to Eq. (2.173), a trivial fundamental group also suggests that the holonomy group stores purely geometrical information. The second extreme case is realised when $\text{Hol}_u^0(\Phi) = \{e\}$, i.e. when the restricted holonomy group of the bundle is trivial. In this case, Eq. (2.173) tells us that for $\text{Hol}_u(\Phi)$ to be non-trivial, we *require* a non-trivial fundamental group $\pi_1(B)$ of the base manifold. After all, $\pi_1(B) = \{e\}$ implies $\text{Hol}_u(\Phi) = \text{Hol}_u^0(\Phi)$ so if additionally $\text{Hol}_u^0(\Phi) = \{e\}$ we immediately have $\text{Hol}_u(\Phi) = \{e\}$. The size of the holonomy group is thus determined by the size of the fundamental group in this scenario. As a consequence, a trivial restricted holonomy group suggests that the holonomy group contains purely topological information. It seems that situations with a trivial fundamental group and those with a trivial restricted holonomy group represent extremes where holonomy reflects only geometric and topological information, respectively. Accordingly, every other combination of a fundamental group and a restricted holonomy group will interpolate between these two extremes. While it appears logical that the size of the fundamental group $\pi_1(B)$ should control the amount of topological information in the holonomy group, it may be less clear why the size of the restricted holonomy group $\text{Hol}_u^0(\Phi)$ should do the same for geometric information. To illustrate why this is the case, we observe that $\text{Hol}_u^0(\Phi)$ captures the local structure of $\text{Hol}_u(\Phi)$ in the sense that all local loops, i.e. all loops within local charts $U \subset B$, are contractible and thus contribute specifically to $\text{Hol}_u^0(\Phi)$. If we understand topology as being concerned with global properties, we find that the only possible “cause” for holonomy on these local charts is of a geometric nature. More precisely, the local origin for holonomy is a geometric property that we call the *curvature* R of the bundle. It is the direct generalisation of the literal curvature of surfaces like the sphere. In fact, the latter corresponds to the example shown in Fig. 2.10, where the concatenation of the red and blue transport curves results in a contractible loop in $\mathbb{S}^2$ that lifts to an open curve in the tangent bundle $T\mathbb{S}^2$ (mismatch between the red and blue arrow at p_1), generating a holonomy that is famously attributed to the curvature of $\mathbb{S}^2$. Since there are many local neighbourhoods $U \simeq \mathbb{R}^n$ in an m -dimensional manifold, the curvature is often regarded as the primary source for holonomy. For this reason, it is often said that the smaller the holonomy group, the flatter the bundle. Indeed, one can show that $\text{Hol}_u^0(\Phi) = \{e\}$ if and only if $R = 0$, i.e. the restricted holonomy group is trivial if and only if the curvature vanishes. In this way, the curvature R serves as a measure for the amount of geometric information in $\text{Hol}_u(\Phi)$ much like the fundamental group serves as a measure for the amount of topological information.

2.2.6 Curvature

In the previous section we introduced curvature as the source of local holonomy and local holonomy as a measure for the non-commutativity of parallel transport. A rigorous definition of curvature therefore has to formalise the idea that *curvature is the failure of parallel transport to commute*. To make this precise, consider a smooth fibre bundle $E \xrightarrow{\pi} B$ with fibre F and an Ehresmann connection HE . The notion of parallel transport under HE is based on the horizontal lift of curves from the base manifold B to the total space E . The lift of a curve is said to be horizontal if its tangent vector is horizontal at all times. Therefore, a horizontal curve $\tilde{\gamma}$ can be regarded as an integral curve $\frac{d}{ds}\tilde{\gamma}(s) = X(\tilde{\gamma}(s))$ of a horizontal vector field X by definition. Here, a horizontal vector field denotes a vector field that is a (local) section $X \in \Gamma(HE, E)$ of the horizontal subbundle HE over E . In this way, the parallel transport of an element $u \in E$ is naturally associated to a horizontal vector field. If we consider the parallel transport along two distinct lifted curves, we get two distinct horizontal vector fields $X, Y \in \Gamma(HE, E)$ and it becomes natural to ask about the flow of either vector field along the other. In particular, do horizontal vector fields remain horizontal if they flow along other horizontal vector fields? If they stay horizontal, the horizontal subspaces at different $u, w \in E$ are said to be compatible and there is no local holonomy.¹⁷ If they fail to stay horizontal, the horizontal subspaces at different $u, w \in E$ are not compatible and the horizontal flow of the local horizontal subspaces develops a vertical component that generates holonomy effects. Determining whether horizontal vector fields remain horizontal under horizontal flow reduces to

¹⁷In this case HE is called integrable because it can be integrated into TE in the sense that we can find a smooth $\dim_{\mathbb{R}}(B)$ -dimensional submanifold I of TE that has the horizontal subbundle for a tangent bundle, i.e. $TI = HE$. Moving around in I corresponds to traversing the aforementioned “horizontal layer” of equivalent points in different fibres of E , which leaves no room for holonomy effects to take place.

an analysis of their closure under the *Lie bracket*. The Lie bracket $[\cdot, \cdot]$ is an operator that assigns to any two vector fields $X, Y : M \rightarrow TM$ on a smooth manifold M another vector field $[X, Y] : M \rightarrow TM$ that is defined via its action

$$[X, Y](f) = X(Y(f)) - Y(X(f)) \quad (2.174)$$

on a smooth function $f \in C^\infty(M)$ on M . Here, the action of a vector field $X \in \Gamma(M)$ on a smooth function $f \in C^\infty(M)$ gives another smooth function $X(f)(\cdot) \equiv X^\mu \frac{\partial f}{\partial x^\mu}(\cdot)$ where x^μ denote local coordinates. Importantly, the Lie bracket $[X, Y]$ is the derivative of Y along the flow generated by X , i.e. the infinitesimal flow of Y along X . The closure of the horizontal vector fields under the Lie bracket is then measured by the expression

$$R(X, Y) := [\Psi(X), \Psi(Y)] - \Psi([\Psi(X), \Psi(Y)]), \quad (2.175)$$

where $[\cdot, \cdot]$ is the Lie bracket and $\Psi := (\text{id} - \Phi) : TE \rightarrow HE$ denotes the projection map onto the horizontal subbundle, defined via the aforementioned Ehresmann projection $\Phi : TE \rightarrow VE$ onto the vertical subbundle. The operator R takes any two vector fields $X, Y : E \rightarrow TE$ on E , projects them onto their horizontal components $\Psi(X), \Psi(Y) : E \rightarrow HE$, and returns the vertical part of the infinitesimal flow $[\Psi(X), \Psi(Y)]$ of $\Psi(Y)$ along $\Psi(X)$ as the discrepancy between the Lie bracket $[\Psi(X), \Psi(Y)]$ and its horizontal part $\Psi([\Psi(X), \Psi(Y)])$. We define the curvature of an Ehresmann connection as the extent of horizontal non-closure that is captured by R .

Definition 2.2.12. Ehresmann Curvature. Let $E \xrightarrow{\pi} B$ be a smooth fibre bundle with an Ehresmann connection HE and let Φ denote the vertical projection map of HE . We denote the complementary horizontal projection map of Φ by $\Psi = (\text{id} - \Phi)$. The **curvature** R of Φ is defined as the vertical component

$$\begin{aligned} R(X, Y) &= [\Psi(X), \Psi(Y)] - \Psi([\Psi(X), \Psi(Y)]) \\ &= \Phi([\Psi(X), \Psi(Y)]) \\ &\equiv [X_H, Y_H]_V \end{aligned} \quad (2.176)$$

of the Lie bracket between the horizontal vector field projections $X_H, Y_H : E \rightarrow HE$ of any two vector fields $X, Y : E \rightarrow TE$.

In the second line of Eq. (2.176) we plugged in $\Psi = (\text{id} - \Phi)$ to arrive at an expression for R that only depends on the Ehresmann projection map Φ . The last line of Eq. (2.176) introduces a shorthand notation of R that will be useful later on. Previously, we showed that Φ can be regarded as a TE -valued one-form on E . As a consequence, R can be understood as a TE -valued *two-form* on E .

In a principal G -bundle $E = P(B, G)$, the Ehresmann connection form Φ defines a connection one-form ω with values in the Lie algebra $\mathfrak{g}$ of G . If we apply the above definition Eq. (2.176) of the curvature to ω instead of Φ , we get a $\mathfrak{g}$ -valued two-form Ω that is called the curvature two-form of the connection one-form ω . To do this, we express $R(X, Y)$ as

$$R(X, Y) = \Phi([X_H, Y_H]), \quad (2.177)$$

where we took the last line of Eq. (2.176) and wrote out the vertical part $[\cdot, \cdot]_V$ of the Lie bracket as $[\cdot, \cdot]_V = \Phi([\cdot, \cdot])$ using the original Ehresmann projection map Φ . Substituting Φ and R with ω and Ω then yields

$$\Omega(X, Y) = \omega([X_H, Y_H]). \quad (2.178)$$

Since ω implements the projection on the vertical subspace, we have $\omega(X_H) = 0$ for every horizontal vector field X_H and we can add two zero terms to arrive at

$$\Omega(X, Y) = \omega([X_H, Y_H]) - X_H \omega(Y_H) + Y_H \omega(X_H), \quad (2.179)$$

which, using the invariant formula

$$d\eta(X, Y) = X\eta(Y) - Y\eta(X) - \eta([X, Y]) \quad (2.180)$$

for the exterior derivative $d\eta$ of a one-form η , becomes

$$\Omega(X, Y) = -d_E\omega(X_H, Y_H), \quad (2.181)$$

where d_E denotes the exterior derivative on the total space E . Note that $d_E\omega(X_H, Y_H)$ is precisely the definition of the covariant derivative

$$D_\omega\omega(X, Y) = d_E\omega(X_H, Y_H) \quad (2.182)$$

of ω capturing its horizontal variation. This means that $\Omega(X, Y)$ can be brought into the compact form $\Omega(X, Y) = -D_\omega\omega(X, Y)$, where the minus sign is often absorbed to arrive at the definition

$$\Omega(X, Y) := D_\omega\omega(X, Y). \quad (2.183)$$

Although this definition of Ω is very elegant, it is not quite as useful for practical applications. To get a more functional expression for Ω , we note that the exterior derivative $d_E\omega(X_H, Y_H)$ of the horizontal vector field projections X_H, Y_H of any two vector fields X, Y can be rewritten as

$$d_E(X_H, Y_H) = d_E\omega(X, Y) + [\omega(X), \omega(Y)], \quad (2.184)$$

such that we arrive at the expression

$$\Omega(X, Y) = d_E\omega(X, Y) + [\omega(X), \omega(Y)] \quad (2.185)$$

for the $\mathfrak{g}$ -valued two-form Ω on E . Finally, one can show that Eq. (2.185) is equivalent to the familiar definition

$$\Omega = d_E\omega + \omega \wedge \omega \quad (2.186)$$

of the curvature two-form Ω in terms of the exterior derivative $d_E\omega$ of ω and the wedge product of connection one-form ω with itself. The advantage of the formulas Eq. (2.185) and Eq. (2.186) is that they only require explicit knowledge of the connection one-form ω since they no longer contain horizontal projections of vector fields. Recall that in Def. 2.2.9 we defined the local connection form $\mathcal{A}_i$ as the pullback $\mathcal{A}_i = \sigma_i^*(\omega)$ of the $\mathfrak{g}$ -valued connection one-form ω by a local section σ_i . The local curvature form $\mathcal{F}_i$ of the $\mathfrak{g}$ -valued two-form Ω is defined analogously.

Definition 2.2.13. Local Curvature Form of a Principal Bundle. Let $E = P(B, G)$ be a principal G -bundle with a $\mathfrak{g}$ -valued Ehresmann connection one-form ω and let $\Omega = d_E\omega + \omega \wedge \omega$ be its $\mathfrak{g}$ -valued curvature two-form. The **local curvature form** $\mathcal{F}_i$ is defined as the pullback

$$\mathcal{F}_i = \sigma_i^*(\Omega) \quad (2.187)$$

of Ω by a local section $\sigma_i : U_i \rightarrow E$ of E .

Specifically, the pullback of Ω from Eq. (2.186) by a local section $\sigma_i : U_i \rightarrow E$ becomes

$$\mathcal{F}_i = d_B\mathcal{A}_i + \mathcal{A}_i \wedge \mathcal{A}_i, \quad (2.188)$$

where $\mathcal{A}_i$ is the local $\mathfrak{g}$ -valued connection one-form and where d_B denotes the exterior derivative on the base space B . The action of $\mathcal{F}_i$ on vectors fields $X, Y : B \rightarrow TB$ on the base manifold B is accordingly given by

$$\mathcal{F}_i(X, Y) = d_B\mathcal{A}_i(X, Y) + [\mathcal{A}_i(X), \mathcal{A}_i(Y)]. \quad (2.189)$$

If we consider a local chart $U_i \subset B$ with coordinates x^μ , we can write the local connection one-form $\mathcal{A}_i$ in components

$$\mathcal{A}_i \equiv \mathcal{A} = \mathcal{A}_\mu dx^\mu, \quad (2.190)$$

where we dropped the subscript i for better readability. From this, we get the component expression

$$\begin{aligned}
\mathcal{F}_i \equiv \mathcal{F} &= d_B \mathcal{A} + \mathcal{A} \wedge \mathcal{A} \\
&= \frac{\partial \mathcal{A}_\nu}{\partial x^\mu} dx^\mu \wedge dx^\nu + \mathcal{A}_\mu \mathcal{A}_\nu dx^\mu \wedge dx^\nu \\
&= \frac{1}{2} \left(\partial_\mu \mathcal{A}_\nu - \partial_\nu \mathcal{A}_\mu + [\mathcal{A}_\mu, \mathcal{A}_\nu] \right) dx^\mu \wedge dx^\nu \\
&=: \mathcal{F}_{\mu\nu} dx^\mu \wedge dx^\nu
\end{aligned} \tag{2.191}$$

of the local curvature form $\mathcal{F}$ on $U \subset B$. Note that in the third line of Eq. (2.191) we wrote $\partial_\mu = \frac{\partial}{\partial x^\mu}$ and used the antisymmetry of the wedge product to explicitly antisymmetrise the components $\mathcal{F}_{\mu\nu}$. On overlapping charts $U_i \cap U_j \neq \emptyset$ the field strengths $\mathcal{F}_i$ and $\mathcal{F}_j$ have to be related as

$$\mathcal{F}_j = t_{ij}^{-1} \mathcal{F}_i t_{ij}, \tag{2.192}$$

where $t_{ij} : U_i \cap U_j \rightarrow G$ are the transition functions between the sections σ_i over U_i and σ_j over U_j . If we compare the transformation condition of the curvature two-form in Eq. (2.192) to that of the connection one-form in Eq. (2.154), we find that for an Abelian gauge group G , only the connection transforms non-trivially while the curvature remains invariant.

Recall that the connection one-form $\mathcal{A}_i$ is identified with the gauge potential in gauge theory. Accordingly, the curvature two-form $\mathcal{F}_i$ corresponds to the field strength, cf. Tab. 2.2. Equation (2.192) then describes the covariant nature of the field strength under gauge transformations. The fact that Eq. (2.192) becomes trivial for Abelian gauge groups tells us that the field strength becomes gauge invariant in such cases. An important class of principal G -bundles to which this applies is the class of principal $U(1)$ -bundles that naturally appear in the topological description of many quantum theories due to the $U(1)$ degree of freedom in the definition of a quantum state.

2.3 Characteristic Classes and Numbers

Characteristic classes are cohomology classes that are naturally associated to vector bundles. As such, they are topological invariants measuring the extent to which the global structure of the bundle deviates from a product structure. Throughout this section we closely follow Ref. [15].

Definition 2.3.1. Characteristic Classes. A **characteristic class** x is a natural assignment of a cohomology class $x(E) \in H^*(B)^{18}$ to each vector bundle $E \xrightarrow{\pi} B$. A characteristic class $x(E)$ of a bundle E is called **stable** if it is invariant under taking the direct sum with a trivial vector bundle T , i.e. if $x(E \oplus T) = x(E)$.

Note that the naturality of the assignment means that the characteristic classes x satisfy

$$x(f^*E) = f^*(x(E)) \quad (2.193)$$

for every continuous map $f : M \rightarrow B$, i.e. that they commute with the notion of pullback. This has various consequences. An important one is that it allows us to compute the characteristic classes of all real and complex vector bundles in terms of the characteristic classes of their classifying spaces $BO(n) = G_n(\mathbb{R}^\infty)$ and $BU(n) = G_n(\mathbb{C}^\infty)$. Following Ref. [15], one may identify four *main types* of characteristic classes:

1. The **Euler class** $e(E) \in H^n(B, \mathbb{Z})$ for oriented, rank- n , real vector bundles $E \xrightarrow{\pi} B$.
2. The **Stiefel–Whitney classes** $w_i(E) \in H^i(B, \mathbb{Z}_2)$ for real vector bundles $E \xrightarrow{\pi} B$.
3. The **Chern classes** $c_i(E) \in H^{2i}(B, \mathbb{Z})$ for complex vector bundles $E \xrightarrow{\pi} B$.
4. The **Pontryagin classes** $p_i(E) \in H^{4i}(B, \mathbb{Z})$ for real vector bundles $E \xrightarrow{\pi} B$.

Except for the Euler class, these classes are universal in that they generate the cohomology rings $H^*(BO(n), \mathbb{Z}_2)$, $H^*(BU(n), \mathbb{Z})$ and $H^*(BSO(n), \mathbb{Z})$ of the respective classifying spaces. We will further address this later on. First, some comments are in order. Although these classes are *characteristic of the vector bundles* $E \xrightarrow{\pi} B$, they correspond to *cohomology classes of the respective base manifolds* B . Furthermore, only the Stiefel–Whitney classes are defined with $\mathbb{Z}_2$ coefficients and only the Chern classes are defined for complex vector bundles. The other classes are defined with $\mathbb{Z}$ coefficients and for real vector bundles. It is also worth mentioning that there is only one Euler class associated to a real vector bundle. In contrast, the other classes come as sequences of $i = 0, 1, \dots, N$ (possibly) non-trivial classes, where the maximal non-trivial index N depends on the specific characteristic class under consideration.¹⁹

One of the most notable features of characteristic classes is that, despite their abstract definition, they often hold remarkably practical meaning. Take for example a real vector bundle $E \xrightarrow{\pi} B$. A first check of the non-triviality of E could be to analyse whether it is orientable or not. One way to do this is by means of a homomorphism $O_E : \pi_1(B) \rightarrow \mathbb{Z}_2$ that assigns 0 or 1 to each loop $\gamma \in \pi_1(B)$ depending on whether the orientation of the fibre is preserved or reversed as one goes around γ . Under the Abelianisation of the fundamental group $\pi_1(B)$ to the first homology group $H_1(B)$, the homomorphism O_E becomes a map $O_E : H_1(B) \rightarrow \mathbb{Z}_2$, which readily defines an element of the first cohomology group $H^1(B, \mathbb{Z}_2)$ with coefficients in $\mathbb{Z}_2$. That element is precisely the first Stiefel–Whitney class, i.e. $[O_E] = w_1(E)$. By construction, $w_1(E)$ is equal to zero if and only if E is orientable. Thus, *non-triviality* of $w_1(E)$ *prevents* E from being orientable. For this reason we call $w_1(E)$ an obstruction (class) to orientability. Most characteristic classes allow an interpretation in terms of obstructions to certain geometric structures. Another striking example of the utility of characteristic classes is the existence of *characteristic numbers* that can be defined from the characteristic classes as follows. Consider a vector bundle $E \xrightarrow{\pi} B$ over an n -dimensional base manifold B . If B is closed and oriented, then there exists a unique homology class $[B] \in H_n(B)$ called the fundamental or orientation class of B that can be naturally paired against elements $\omega \in H^n(B, R)$ of the n -th cohomology group with coefficients in R to give an element of the coefficient ring R of $H^n(B, R)$, i.e.

$$\langle [\omega], [B] \rangle \equiv \omega(B) \in R. \quad (2.194)$$

¹⁸ $H^*(B)$ denotes the cohomology ring of the base manifold B as defined in Eq. (2.83).

¹⁹The cohomology groups $H^n(B)$ of the base manifold B become trivial once $n > \dim_{\mathbb{R}}(B)$. The maximal index of the Stiefel–Whitney/Chern/Pontryagin classes is therefore the maximal number N such that $N/2N/4N \leq \dim_{\mathbb{R}}(B)$.

In a vector bundle with characteristic classes $\{x_i\}$, the cup product of cohomology then allows us compute products $x_{i_1}x_{i_2}\dots x_{i_p} := x_{i_1} \smile x_{i_2} \smile \dots \smile x_{i_p}$ of characteristic classes of total degree $\sum_{k=1}^p \deg(x_{i_k}) = n$ and pair them against the orientation class $[B]$ to get an element $(x_{i_1}x_{i_2}\dots x_{i_p})(B) \in R$ of the coefficient ring that is called a characteristic number and that may acquire a geometric or even physical meaning in some cases.²⁰

Definition 2.3.2. Characteristic Numbers. A **characteristic number** X of a vector bundle $E \xrightarrow{\pi} B$ over an oriented closed n -dimensional manifold B is the result of any pairing

$$X := (x_{i_1}\dots x_{i_p})(B) \quad (2.195)$$

between a cup product $(x_{i_1}\dots x_{i_p})$ of characteristic classes x_{i_k} of total degree $\sum_{k=1}^p \deg(x_{i_k}) = n$ and the orientation class $B \equiv [B] \in H_n(B)$ of B .

More generally, one can pair an element $\omega \in H^n(B, R)$ of the n -th cohomology group with coefficients in R against any orientable n -cycle $C \subset B$ of dimension $\dim_{\mathbb{R}}(C) = n \leq \dim_{\mathbb{R}}(B)$ to get an element $\langle [\omega], [C] \rangle \in R$. For $C \neq B$, the resulting ring elements are simply not considered “characteristic” of the bundle, as they depend not only on the bundle but on the specific connection and sub-cycle $C \subset B$.

In the following, we will give a brief overview over some of the characteristic classes most relevant to physics, namely the Euler, Stiefel–Whitney, and Chern classes.

2.3.1 The Euler Class

In 1758, Euler famously discovered that the alternating sum of the number of vertices V , edges E and faces F of a convex polyhedron $P \subset \mathbb{R}^3$ always gives

$$V - E + F = 2. \quad (2.196)$$

This equation is now known as Euler’s polyhedron formula and the sum of vertices, edges and faces is called the Euler characteristic $\chi(P)$ of P . A modern version of the Euler characteristic generalises this surprising result to general topological spaces: the Euler characteristic $\chi(X)$ of an n -dimensional topological space X is defined as the alternating sum

$$\chi(X) = \sum_{i=0}^n (-1)^i \beta_i \quad (2.197)$$

of its Betti numbers $\beta_i = \text{rank}(H_i(X))$, making it a topological homotopy invariant of X . Since every convex polyhedron is homotopic to the two-sphere, Eq. (2.196) is a special case of Eq. (2.197) for $X = \mathbb{S}^2$. Indeed, we can read off the Betti numbers $\beta_0 = \beta_2 = 1$ and $\beta_1 = 0$ in Eq. (2.56) and find

$$\chi(\mathbb{S}^2) = \beta_0 - \beta_1 + \beta_2 = 1 - 0 + 1 = 2. \quad (2.198)$$

Based on Eq. (2.197), we can extend this formula to higher-dimensional spheres. For example, with Eq. (2.56) the Euler characteristic of the three-dimensional sphere $\mathbb{S}^3$ becomes

$$\chi(\mathbb{S}^3) = \beta_0 - \beta_1 + \beta_2 - \beta_3 = 1 - 0 + 0 - 1 = 0, \quad (2.199)$$

and we arrive at the closed expression

$$\chi(\mathbb{S}^d) = 1 + (-1)^d \quad (2.200)$$

for the Euler characteristic of the d -sphere. Historically, the Euler characteristic was the first example of a characteristic number. It can be used to define the Euler class $e(E)$ of a real oriented rank- n vector bundle $E \xrightarrow{\pi} B$ over an oriented closed manifold B as the unique cohomology class $e(E) \in H^n(B, \mathbb{Z})$ that satisfies

$$\langle e(E), [B] \rangle = \chi(E), \quad (2.201)$$

²⁰There are subtleties when products involve characteristic classes from cohomology groups with different coefficients.

i.e. that results in the Euler characteristic $\chi(E)$ of the bundle $E \xrightarrow{\pi} B$ when paired with the orientation class $[B] \in H_n(B)$ of the base manifold B .²¹ Equation (2.201) immediately implies that the Euler class $e(E)$ changes sign when the opposite orientation is chosen for E , i.e. that

$$e(\bar{E}) = -e(E), \quad (2.202)$$

where $\bar{E}$ is E with reversed orientation. This has an important consequence for the Euler class of real oriented vector bundles of odd rank, in which the fibre inversion isomorphism $I : F \rightarrow F, f \mapsto -f$ reverses the orientation²² of the fibre so that it induces a vector bundle isomorphism $\bar{I} : E \rightarrow \bar{E}$ between $\bar{E}$ and E . It follows, that the Euler class of real oriented vector bundles of odd rank satisfies

$$e(E) \stackrel{\bar{I}}{=} e(\bar{E}) = -e(E), \quad (2.203)$$

which leads to the equation

$$e(E) + e(E) = 2e(E) = 0. \quad (2.204)$$

Note that Eq. (2.204) does not necessarily imply that $e(E) = 0$ since the cohomology group $H^n(B, \mathbb{Z})$ may contain torsion elements, i.e. elements $g \in H^n(B, \mathbb{Z})$ for which there exists a finite integer m such that $m \cdot g = 0$. Therefore, Eq. (2.204) only tells us that the Euler class of real oriented vector bundles of odd rank is 2-torsion. This has some implications for the Chern–Weil perspective on the Euler class which we will address in greater detail shortly.

Unlike most other characteristic classes, the Euler class is *unstable*, meaning that $e(E \oplus T) \neq e(E)$ for the direct sum $E \oplus T$ between a given vector bundle E and a trivial vector bundle T . In fact, it satisfies

$$e(E \oplus T) = e(E) \smile e(T) = 0. \quad (2.205)$$

Note that the instability of the Euler class is by no means a flaw. In fact, it makes the Euler class one of the few characteristic classes that have any chance of detecting the non-triviality of so-called stably trivial vector bundles, i.e. vector bundles E for which there exists a trivial vector bundle T_E such that the direct sum $E \oplus T_E$ is again a trivial bundle T , i.e. $E \oplus T_E = T$. Stable characteristic classes of stably trivial vector bundles are therefore necessarily trivial themselves, as $x(E) = x(E \oplus T_E) = x(T) = 0$.

Conceptually, the Euler class is an obstruction to finding a nowhere vanishing global section on a vector bundle. This means that $e(E) = 0$ is a necessary condition for the existence of a nowhere vanishing global section. Importantly, $e(E) = 0$ is not sufficient – there may be situations where $e(E) = 0$ and no nowhere vanishing global sections exist. Recall that a rank- n vector bundle E is trivial if and only if it admits a global frame, i.e. a collection of n linearly independent (orthonormal) nowhere vanishing sections. The Euler class is therefore an obstruction to finding a trivial rank-one subbundle of E . From this perspective, Eq. (2.200) tells us that there exist no nowhere vanishing global sections on the tangent bundles of even-dimensional spheres, which is precisely the statement of the hairy ball theorem. The modern textbook definition of the Euler class formalises the idea of obstructing nowhere vanishing sections by means of the so-called Thom isomorphism [61]. This isomorphism is used to translate the twisting of the zero-section E_0 within $E \xrightarrow{\pi} B$ into a concrete cohomology class $e(E) \in H^n(B)$ by relating the orientation class of the total bundle E to the complement $E - E_0$ of E_0 in E , i.e. the domain of nowhere vanishing sections of E .

²¹The Euler characteristic $\chi(M)$ of a smooth manifold M is the Euler class $e(TM)$ of its tangent bundle TM .

²²The fibre inversion is implemented by the negative identity matrix $\mathbb{1}_n$ which has determinant $\det(\mathbb{1}_n) = (-1)^n$ which equals plus one (preserves orientation) for even n and minus one (reverses orientation) for odd n .

2.3.2 The Stiefel–Whitney Classes

Just like the Euler class, the Stiefel–Whitney classes are defined for real vector bundles $E \xrightarrow{\pi} B$. However, the Stiefel–Whitney classes do not require orientability because they are defined in cohomology with $\mathbb{Z}_2$ coefficients. The following list of properties of the Stiefel–Whitney classes can be regarded as their axiomatic definition [15].

Definition 2.3.3. Stiefel–Whitney Classes. The **Stiefel–Whitney classes** are a unique non-trivial sequence of functions $\{w_i\}$ that assign cohomology classes $w_i \in H^i(B, \mathbb{Z}_2)$ to every real vector bundle $E \xrightarrow{\pi} B$ with typical fibre F such that

1. $w_i(f^*E) = f^*(w_i(E))$ for every continuous function $f : M \rightarrow B$.
2. $w(E_1 \oplus E_2) = w(E_1) \smile w(E_2)$ where $w = \sum_i w_i \in H^*(B, \mathbb{Z}_2)$.
3. $w_i(E) = 0$ if $i > \dim_{\mathbb{R}}(F)$ or $i > \dim_{\mathbb{R}}(B)$.

The sum $w = \sum_i w_i$ is called the total Stiefel–Whitney class of E .

The first condition in the above list makes the Stiefel–Whitney classes natural in the aforementioned sense. The second condition yields the relation

$$w_n(E_1 \oplus E_2) = \sum_{i+j=n} w_i(E_1) \smile w_j(E_2), \quad (2.206)$$

which is sometimes called the Whitney sum formula. Equation (2.206) readily establishes the stability of the Stiefel–Whitney classes: the n -th Stiefel–Whitney class $w_n(E \oplus T)$ of the direct sum $E \oplus T$ of a given vector bundle E and a trivial vector bundle T is

$$w_n(E \oplus T) = \sum_{i+j=n} w_i(E_1) \smile w_j(E_2) = w_n(E) \smile w_0(T) = w_n(E) \smile 1 = w_n(E), \quad (2.207)$$

since $w_j(T) = 0$ for all $j > 0$. The third condition of Def. 2.3.3 ensures that only finitely many terms of the total Stiefel–Whitney class are non-zero. The Stiefel–Whitney classes do not require orientability. Nonetheless, there is a connection to the Euler class that is unique for orientable vector bundles.

Theorem 2.3.1. *Let $E \xrightarrow{\pi} B$ be an orientable rank- n real vector bundle. The top Stiefel–Whitney class $w_n(E)$ and the Euler class $e(E)$ of E are related as*

$$e(E) = w_n(E) \pmod{2}. \quad (2.208)$$

This relationship is not accidental. One way to construct the Stiefel–Whitney classes is to take Thm. 2.3.1 as a starting point and define the other Stiefel–Whitney classes by induction. This construction scheme draws on the history of characteristic classes, telling a story about the stabilisation and extension of the Euler class to non-orientable spaces.

We mentioned earlier that the first Stiefel–Whitney class $w_1(E)$ is an obstruction to the orientability of a given vector bundle E . The second Stiefel–Whitney class provides a refinement of this notion that is of particular interest to physicists: it obstructs the existence of spin structures on a given vector bundle. To see how this is similar to an obstruction to orientability, take a generic real rank- n vector bundle E with the general linear group $\text{GL}(n)$ for a structure group. If E is equipped with a metric we can reduce $\text{GL}(n)$ to the orthogonal group $\text{O}(n)$ which preserves the metric – this is a typical starting point for physical theories. Orientability of E then allows us to further reduce $\text{O}(n)$ to the special orthogonal group $\text{SO}(n)$ which preserves orientation. Finally, E is said to admit a spin structure if the structure group $\text{SO}(n)$ can be lifted to its double cover, the spin group $\text{Spin}(n)$, in a consistent fashion. While the first Stiefel–Whitney class $w_1(E)$ obstructs orientability, i.e. the reduction to $\text{SO}(n)$, the second Stiefel–Whitney class $w_2(E)$ further obstructs the lift of $\text{SO}(n)$ to $\text{Spin}(n)$. Therefore, a real vector bundle $E \xrightarrow{\pi} B$ admits a spin structure if and only if $w_1(E) = w_2(E) = 0$ [15].

Finally, the Stiefel–Whitney classes are universal classes of real vector bundles. Specifically, the Stiefel–Whitney classes generate the cohomology ring $H^*(BO(n), \mathbb{Z}_2)$ of the aforementioned classifying space $BO(n) = G_n(\mathbb{R}^\infty)$ of all real rank- n vector bundles and associated principal $\text{O}(n)$ -bundles. This means that every characteristic class of a real vector bundle that is defined with $\mathbb{Z}_2$ coefficients is a polynomial in the Stiefel–Whitney classes. The reason we care most about cohomology classes with $\mathbb{Z}_2$ coefficients for real vector bundles is that real vector bundles, unlike their complex counterparts, may not be orientable and $\mathbb{Z}_2$ coefficients work for both orientable and non-orientable spaces equally.

2.3.3 The Chern Classes

The Chern classes are defined for complex vector bundles $E \xrightarrow{\pi} B$, i.e. for vector bundles whose typical fibre F is a complex vector space. Accordingly, the structure group of a generic complex rank- n vector bundle is the complex general linear group $\mathrm{GL}(n, \mathbb{C})$, which, unlike the real general linear group $\mathrm{GL}(n, \mathbb{R})$, is a *connected* topological manifold. One immediate consequence of this is that every complex vector bundle admits a canonical orientation. In contrast to the Euler class and the Stiefel–Whitney classes, the Chern classes therefore do not require any mention of orientability and can be defined as cohomology classes with $\mathbb{Z}$ coefficients, extending the Euler class to complex spaces. The following list of properties of the Chern classes can be regarded as their axiomatic definition [15].

Definition 2.3.4. Chern Classes. The **Chern classes** are a unique non-trivial sequence of functions $\{c_i\}$ that assign cohomology classes $c_i \in H^{2i}(B, \mathbb{Z})$ to every complex vector bundle $E \xrightarrow{\pi} B$ with typical fibre F such that

1. $c_i(f^*E) = f^*(c_i(E))$ for every continuous function $f : M \rightarrow B$.
2. $c(E_1 \oplus E_2) = c(E_1) \smile c(E_2)$ where $c = \sum_i c_i \in H^*(B, \mathbb{Z})$.
3. $c_i(E) = 0$ if $i > \dim_{\mathbb{C}}(F)$ or $2i > \dim_{\mathbb{R}}(B)$.

The sum $c = \sum_i c_i$ is called the total Chern class of E .

Again, the first condition in the above list makes the Chern classes natural in the aforementioned sense and again the second condition yields the Whitney sum formula

$$c_n(E_1 \oplus E_2) = \sum_{i+j=n} c_i(E_1) \smile c_j(E_2) \quad (2.209)$$

for the Chern classes, establishing their stability as

$$c_n(E \oplus T) = \sum_{i+j=n} c_i(E) \smile c_j(T) = c_n(E) \smile c_0(T) = c_n(E) \smile 1 = c_n(E), \quad (2.210)$$

since $c_j(T) = 0$ for all $j > 0$. Just like with the Stiefel–Whitney classes, the third condition of Def. 2.3.4 ensures that only finitely many terms of the total Chern class are non-zero.

We mentioned in the beginning that every complex vector bundle E comes with a canonic orientation. This canonic orientation induces an orientation on the rank- $2n$ real vector bundle $E_{\mathbb{R}} \xrightarrow{\pi} B$ that is associated to any rank- n complex vector bundle $E \xrightarrow{\pi} B$ by viewing the typical fibre $F \simeq \mathbb{C}^n$ as a real vector space $F_{\mathbb{R}} \simeq \mathbb{R}^{2n}$. The Chern classes are designed such that there is a strong relation between the top Chern class $c_n(E)$ and the Euler class $e(E_{\mathbb{R}})$.

Theorem 2.3.2. *Let $E \xrightarrow{\pi} B$ be a rank- n complex vector bundle and let $E_{\mathbb{R}} \xrightarrow{\pi} B$ be the orientable rank- $2n$ real vector bundle associated to it. The top Chern class $c_n(E)$ is equal to the Euler class $e(E_{\mathbb{R}})$,*

$$e(E_{\mathbb{R}}) = c_n(E). \quad (2.211)$$

Again, this relationship is not accidental. Just like the Stiefel–Whitney classes, the Chern classes can be defined by taking Thm. 2.3.2 as a starting point and defining the lower Chern classes by induction [72]. This construction scheme of the Chern classes represents the stabilisation and extension of the Euler class to complex spaces.

Theorem (2.3.2) tells us that the top Chern class $c_n(E)$ of a rank- n complex vector bundle E obstructs the existence of a nowhere vanishing global section, i.e. of a trivial complex rank-1 subbundle of E . The lower Chern classes then extend this notion to trivial subbundles of higher and higher rank. Specifically, the obstruction to finding k orthonormal global sections in a complex vector bundle of rank n is measured by the $(n - k + 1)$ -th Chern class $c_{n-k+1}(E)$.

Finally, the Chern classes are universal classes of complex vector bundles: they generate the cohomology ring $H^*(BU(n), \mathbb{Z})$ of the classifying space $BU(n) = G_n(\mathbb{C}^{\infty})$ of all complex rank- n vector bundles and associated principal $U(n)$ -bundles. This means that every characteristic class of a complex vector bundle that is defined with $\mathbb{Z}$ coefficients is a polynomial in the Chern classes. We can afford to focus on the richer cohomology with $\mathbb{Z}$ coefficients here because of the aforementioned canonic orientation of complex vector bundles.

2.3.4 The Pontryagin Classes

The Chern classes and their universality for complex vector bundles can be used to define universal characteristic classes with $\mathbb{Z}$ coefficients for real vector bundles. This is done by complexifying a given real vector bundle $E \xrightarrow{\pi} B$ of real rank n as $E_{\mathbb{C}} := E \otimes \mathbb{C} = E \oplus iE$ to get a complex vector bundle $E_{\mathbb{C}} \xrightarrow{\pi} B$ of *complex* rank n , but *real* rank $2n$. Based on the complexification of E , one can then define the so-called Pontryagin classes

$$p_i(E) := (-1)^i c_{2i}(E_{\mathbb{C}}) \in H^{4i}(B, \mathbb{Z}) \quad (2.212)$$

of the initial real vector bundle E . There are some things worth mentioning about Eq. (2.212). First of all, it defines a characteristic class $p_i(E)$ of a real vector bundle, which may lack orientability, via a characteristic class $c_i(E_{\mathbb{C}})$ of a canonically oriented complex vector bundle. As we will see shortly, this discrepancy in orientability introduces some subtleties to the Pontryagin classes. Furthermore, the relation $p_i(E) \propto c_{2i}(E_{\mathbb{C}}) \in H^{4i}(B, \mathbb{Z})$ and $c_{2i}(E_{\mathbb{C}}) = 0$ for $2i > n$ means that there are at most $\lfloor \frac{n}{2} \rfloor$ Pontryagin classes $p_1, \dots, p_{\lfloor \frac{n}{2} \rfloor}$ associated to a given real rank- n vector bundle E . An important consequence of this is that only real vector bundles E of even rank have a Pontryagin class $p_{\frac{n}{2}}(E)$ that corresponds to the top degree Chern class $c_n(E_{\mathbb{C}})$ of the complexified bundle. By Thm. 2.3.2 we then have

$$p_{\frac{n}{2}}(E) \xrightarrow{(2.212)} c_n(E_{\mathbb{C}}) \xrightarrow{(2.211)} e(E \oplus E), \quad (2.213)$$

where $E \oplus E$ is the oriented real bundle associated to the complexified bundle $E_{\mathbb{C}} = E \oplus iE$ of E . However, the Whitney sum formula $e(E \oplus E) = e(E) \smile e(E) = e(E)^2$ that would relate the highest Pontryagin class $p_{\frac{n}{2}}(E)$ to the Euler class $e(E)$ as

$$p_{\frac{n}{2}}(E) = e(E)^2 \quad (2.214)$$

is only valid if the original real vector bundle E is oriented. The fact that the relation between the highest Pontryagin class and the Euler class depends on the rank and the orientability of a vector bundle suggests that the universality of the Pontryagin classes for real vector bundles requires more careful consideration. Indeed, we explicitly have to distinguish between oriented and unoriented vector bundles of even and odd rank.²³ In practice, this means that we have to consider different classifying spaces, namely $BSO(n) = \tilde{G}_n(\mathbb{R}^{\infty})$ for oriented real rank- n bundles and $BO(n) = G_n(\mathbb{R}^{\infty})$ for unoriented real rank- n bundles.²⁴

In the unoriented case, the Euler class e is not defined. One can show that the Pontryagin classes $[p_1, \dots, p_{\lfloor \frac{n}{2} \rfloor}]$ generate the torsion free part $H_{\text{Free}}^*(BO(n), \mathbb{Z}) = H^*(BO(n), \mathbb{Z}) / \text{torsion}$ of the cohomology ring $H^*(BO(n), \mathbb{Z})$, while torsion is generated by (the Bockstein) images of the Stiefel–Whitney classes in $H^*(BO(n), \mathbb{Z})$. Therefore, the integral cohomology $H^*(BO(n), \mathbb{Z})$ is fully determined by the Pontryagin classes and the Stiefel–Whitney classes. Notably, this is true for all n , so that every integral characteristic class of an unoriented real rank- n vector bundle is a polynomial in the Pontryagin classes and (the Bockstein) images of the Stiefel–Whitney classes in $H^*(BO(n), \mathbb{Z})$ irrespective of the whether n is odd or even.

In the oriented case, the Euler class e is defined and the situation becomes a bit more intricate. The torsion part of $H^*(BSO(n), \mathbb{Z})$ is again captured by (the Bockstein) images of the Stiefel–Whitney classes in $H^*(BSO(n), \mathbb{Z})$. However, the free part $H_{\text{Free}}^*(BSO(n), \mathbb{Z})$ of the cohomology ring $H^*(BSO(n), \mathbb{Z})$ requires a distinction between even and odd ranks. For odd ranks, there is no relation between the Pontryagin classes and the Euler class and $H_{\text{Free}}^*(BSO(2m+1), \mathbb{Z})$ is generated by the Pontryagin classes $[p_1, \dots, p_m]$. For even ranks, the highest Pontryagin class p_m equals the square of the Euler class e as $p_m = e^2$ and $H_{\text{Free}}^*(BSO(2m), \mathbb{Z})$ is instead generated by $[p_1, \dots, p_{m-1}, e]$. Therefore, the integral cohomology $H^*(BSO(n), \mathbb{Z})$ is completely determined by the Pontryagin classes and the Stiefel–Whitney classes for odd ranks, while the only additional characteristic class required for even ranks is the Euler class [15].

²³So far we did not have to explicitly take orientability into account because we were either dealing with $\mathbb{Z}_2$ cohomology which ignores orientation or with complex bundles that are canonically oriented to begin with.

²⁴The classifying space for oriented real rank- n vector bundles is $BSO(n) = \tilde{G}_n(\mathbb{R}^{\infty})$ where $\tilde{G}_n(\mathbb{R}^{\infty}) = \lim_{k \rightarrow \infty} \tilde{G}_n(\mathbb{R}^k)$ denotes the limit of real oriented Grassmannians $\tilde{G}_n(\mathbb{R}^k) = SO(k)/(SO(n) \times SO(k-n))$.

2.3.5 Chern–Weil Theory

There is one approach to the theory of characteristic classes that is distinguished by its practical utility. Chern–Weil theory is based on the insight that certain polynomials in the curvature two-form $\mathcal{F}$ of principal G -bundles $P \xrightarrow{\pi} B$ define elements of the de Rham cohomology ring $H_{\text{dR}}^*(B)$ of their base spaces. In this way, Chern–Weil theory provides a method for computing some characteristic classes of principal G -bundles using only their curvature. In the following, we sketch the fundamental ideas of Chern–Weil theory for principal G -bundles. This part is based on Ref. [39].

As a preliminary consideration, we introduce the notion of invariant polynomials. Let G be a Lie group and let $\mathfrak{g}$ be its Lie algebra. An invariant polynomial of degree k is a symmetric k -linear function $f : \otimes^k \mathfrak{g} \rightarrow \mathbb{R}$ that satisfies

$$f(\text{Ad}_g A_1, \dots, \text{Ad}_g A_k) = f(A_1, \dots, A_k) \quad (2.215)$$

for all $g \in G$ and $A_i \in \mathfrak{g}$. Here, $\text{Ad}_g A_i = g^{-1} A_i g$ is the adjoint action of G on $\mathfrak{g}$. The invariant polynomials of degree k form a vector space $I^k(\mathfrak{g})$ and the direct sum $I^*(\mathfrak{g}) = \bigoplus_{k=0}^{\infty} I^k(\mathfrak{g})$ is called the algebra of invariant polynomials. We can extend the domain of invariant polynomials $I^*(\mathfrak{g})$ from $\mathfrak{g}$ to $\mathfrak{g}$ -valued differential forms by defining

$$f(\eta_1 A_1, \dots, \eta_k A_k) := \eta_1 \wedge \dots \wedge \eta_k f(A_1, \dots, A_k), \quad (2.216)$$

where $A_i \in \mathfrak{g}$ and η_i are k_i -forms. The result of Eq. (2.216) is an $\mathbb{R}$ -valued $k = \sum_i k_i$ form. If we take a particular $\mathfrak{g}$ -valued k -form η and an invariant polynomial $f \in I^p(\mathfrak{g}) \subset I^*(\mathfrak{g})$ of degree p we call the $\mathbb{R}$ -valued kp -form $f(\eta) := f(\eta, \dots, \eta)$ the diagonal combination of f in η .

Consider a principal G -bundle $P \xrightarrow{\pi} B$ whose structure group G is a Lie group with Lie algebra $\mathfrak{g}$. Any connection on P allows us to define a $\mathfrak{g}$ -valued connection one-form $\mathcal{A}$ on B and its $\mathfrak{g}$ -valued curvature two-form $\mathcal{F}$ as shown in Eq. (2.188). The Chern–Weil theorem states the following important result about invariant polynomials in the curvature two-form.

Definition 2.3.5. Chern–Weil Theorem. Let $P \xrightarrow{\pi} B$ be a principal G -bundle where G is a Lie group and $\mathfrak{g}$ its Lie algebra. The diagonal combination $f(\mathcal{F}) := f(\mathcal{F}, \dots, \mathcal{F})$ of any invariant polynomial $f \in I^*(\mathfrak{g})$ in a curvature two-form $\mathcal{F}$ on B satisfies

1. $df(\mathcal{F}) = 0$
2. $f(\mathcal{F}') - f(\mathcal{F})$ is exact for any two curvature two-forms $\mathcal{F}'$ and $\mathcal{F}$ on P

The first condition makes $f(\mathcal{F})$ a closed $\mathbb{R}$ -valued differential form on B , while the second condition ensures that $f(\mathcal{F})$ and $f(\mathcal{F}')$ are cohomologous. Combined, the Chern–Weil theorem asserts, that $[f(\mathcal{F})] \in H_{\text{dR}}^*(B)$, i.e. that $f(\mathcal{F})$ is a well-defined representative of a de Rham cohomology group of the base space B . In this way, the Chern–Weil theorem defines a homomorphism

$$\begin{aligned} \phi : I^*(\mathfrak{g}) &\rightarrow H_{\text{dR}}^*(B) \simeq H^*(B, \mathbb{R}) \\ f &\mapsto [f(\mathcal{F})], \end{aligned} \quad (2.217)$$

that is sometimes called the Chern–Weil homomorphism [39]. Recall that the de Rham theorem establishes an equivalence $H_{\text{dR}}^*(B) \simeq H^*(B, \mathbb{R})$ between de Rham cohomology and singular cohomology with *real* coefficients. In contrast, all of the characteristic classes we discussed before are defined with $\mathbb{Z}$ or $\mathbb{Z}_2$ coefficients. The Chern–Weil homomorphism becomes applicable to the integral characteristic classes because the inclusion map $\iota : \mathbb{Z} \hookrightarrow \mathbb{R}$ naturally induces a map $\iota^* : H^*(B, \mathbb{Z}) \rightarrow H^*(B, \mathbb{R})$ between integral and real cohomology. In terms of characteristic classes, the Chern–Weil homomorphism is often stated as follows.

Definition 2.3.6. Chern–Weil Homomorphism. Consider a principal G -bundle $P \xrightarrow{\pi} B$ whose structure group is a Lie group G with Lie algebra $\mathfrak{g}$. Let $\mathcal{F}$ denote its $\mathfrak{g}$ -valued curvature two-form on B and let $x_{\mathbb{R}}(P) \in H^*(B, \mathbb{R})$ denote the image of an integral characteristic class $x(P) \in H^*(B, \mathbb{Z})$ of P , induced by the inclusion $\iota : \mathbb{Z} \hookrightarrow \mathbb{R}$. Then there exists an invariant polynomial $f \in I^*(\mathfrak{g})$ and a homomorphism $\phi : I^*(\mathfrak{g}) \rightarrow H^*(B, \mathbb{R})$ such that

$$\phi(f) = [f(\mathcal{F})] = x_{\mathbb{R}}(P). \quad (2.218)$$

In the following, we briefly address the invariant polynomials that give rise to the Chern classes, the Pontryagin classes, and the Euler class.

Consider a complex vector bundle $E \xrightarrow{\pi} B$ with typical fibre $F \simeq \mathbb{C}^n$. The structure group G of E is $U(n)$, and every (local) connection one-form $\mathcal{A}$ and curvature two-form $\mathcal{F}$ take their values in the Lie algebra $\mathfrak{g} = \mathfrak{u}(n)$. The total Chern class is defined via the invariant polynomial [39]

$$c(\mathcal{F}) := \det \left(\mathbb{1} + \frac{i\mathcal{F}}{2\pi} \right). \quad (2.219)$$

Since $\mathcal{F}$ is a two-form, $c(\mathcal{F})$ is direct sum of forms of even degree, i.e.

$$c(\mathcal{F}) = 1 + c_1(\mathcal{F}) + c_2(\mathcal{F}) + \dots, \quad (2.220)$$

where the $[c_i(\mathcal{F})] \in H^{2i}(B, \mathbb{R})$ are the images of the integral Chern classes $c_i(E) \in H^{2i}(B, \mathbb{Z})$ induced by the inclusion $\iota : \mathbb{Z} \hookrightarrow \mathbb{R}$.

Similarly, let $E \xrightarrow{\pi} B$ be a real vector bundle with typical fibre $F \simeq \mathbb{R}^n$. The structure group G of E is $O(n)$, and every (local) connection one-form $\mathcal{A}$ and curvature two-form $\mathcal{F}$ take their values in the Lie algebra $\mathfrak{g} = \mathfrak{o}(n)$. According to Eq. (2.212), the Pontryagin classes of E are defined via the Chern classes of its complexification $E_{\mathbb{C}}$. This is captured by the invariant polynomial

$$p(\mathcal{F}) := \det \left(\mathbb{1} + \frac{\mathcal{F}}{2\pi} \right) \quad (2.221)$$

capturing the total Pontryagin class. Recall that the generators of $\mathfrak{g} = \mathfrak{o}(n)$ are the skew-symmetric $n \times n$ matrices. Since $\mathcal{F}$ takes its values in $\mathfrak{o}(n)$, it is skew-symmetric as well and we find

$$p(\mathcal{F}) := \det \left(\mathbb{1} + \frac{\mathcal{F}}{2\pi} \right) = \det \left(\mathbb{1}^{\top} + \frac{\mathcal{F}^{\top}}{2\pi} \right) = \det \left(\mathbb{1} - \frac{\mathcal{F}}{2\pi} \right) = p(-\mathcal{F}), \quad (2.222)$$

so $p(\mathcal{F})$ is an *even* polynomial in $\mathcal{F}$. As a result, the polynomials $p_i(E)$ in the expansion

$$p(\mathcal{F}) = 1 + p_1(\mathcal{F}) + p_2(\mathcal{F}) + \dots \quad (2.223)$$

are of order $4i$ and define elements $[p_i(\mathcal{F})] \in H^{4i}(B, \mathbb{R})$ that correspond to the images of the integral Pontryagin classes $p_i(E) \in H^{4i}(B, \mathbb{Z})$ induced by the inclusion $\iota : \mathbb{Z} \hookrightarrow \mathbb{R}$.

The invariant polynomial of the Euler class can be deduced from the Pontryagin classes. Based on Eqs. (2.221) and (2.223) one can show that the highest Pontryagin class $p_{\lfloor \frac{n}{2} \rfloor}$ of a given rank- n vector bundle takes the form

$$p_{\lfloor \frac{n}{2} \rfloor} = \det \left(\frac{\mathcal{F}}{2\pi} \right). \quad (2.224)$$

We explained earlier that for even n , the highest Pontryagin class $p_{\frac{n}{2}}$ fulfills

$$p_{\frac{n}{2}}(E) = e(E)^2, \quad (2.225)$$

which suggests the invariant polynomial

$$e(\mathcal{F}) = \text{Pf} \left(\frac{\mathcal{F}}{2\pi} \right) \quad (2.226)$$

for the Euler class. Indeed, we find that the Pfaffian yields a class $[e(\mathcal{F})] \in H^n(B, \mathbb{R})$ that corresponds to the image of the integral Euler class $e(E) \in H^n(B, \mathbb{Z})$ induced by the inclusion $\iota : \mathbb{Z} \hookrightarrow \mathbb{R}$. Importantly, the inclusion of integral cohomology into real cohomology removes torsion. This means that $e(E)$ is sent to zero whenever it is of finite order. The fact that the Pfaffian of a skew-symmetric matrix of odd size vanishes identically therefore captures the torsion nature of the Euler class in odd rank bundles mentioned in Eq. (2.204).

The characteristic classes $x(\mathcal{F})$ of Chern–Weil theory correspond to de Rham cohomology classes constructed from the curvature two-form. As a result, the characteristic number pairing from Eq. (2.194) can be expressed via the de Rham pairing from Eq. (2.104), yielding the simple formula

$$X(E) := \langle [x(E)], [B] \rangle = \oint_B x(\mathcal{F}) \quad (2.227)$$

for the characteristic numbers of a given vector bundle $E \xrightarrow{\pi} B$.

2.3.6 Chern Numbers and a Glimpse of the Atiyah–Singer Index Theorem

When it comes to topological invariants of vector bundles, characteristic numbers stand out for at least two reasons: their conceptual simplicity and practical utility. The latter extends well beyond pure mathematics. This is impressively illustrated by the integer quantum Hall effect, whose quantised Hall conductance is *directly* determined by a characteristic number known as the first Chern number. This connection was established in 1982 by Thouless, Kohmoto, Nightingale, and den Nijs [10], and proved influential enough to earn Thouless a share of the 2016 Nobel Prize in Physics [73]. Today, Chern numbers are among the most prominent characteristic numbers in theoretical condensed matter physics. In particular, they appear as principal topological invariants in the celebrated tenfold classification of topological quantum matter with symmetries [14]. Considering their ubiquity, it is all the more surprising that an exact definition of Chern numbers remains somewhat ambiguous.

Mathematically, Chern numbers are typically defined as pairings between top-degree products of Chern classes and the orientation class of a base manifold B . Consider, for instance, a complex vector bundle $E \xrightarrow{\pi} B$ over an orientable, closed base manifold B of real dimension $\dim_{\mathbb{R}}(B) = 6$. There are in principle three independent Chern numbers

$$C_1(E) = \langle [c_1^3(E)], [B] \rangle, \quad C_2(E) = \langle [c_1 c_2(E)], [B] \rangle, \quad C_3(E) = \langle [c_3(E)], [B] \rangle \quad (2.228)$$

associated with the three top degree products $c_1^3, c_1 c_2, c_3 \in H^6(B, \mathbb{Z})$ of Chern classes. By definition of the Chern classes, the pairings in Eq. (2.228) (and integer linear combinations thereof) yield integer-valued characteristic numbers of E .

In the physics literature, on the other hand, the term Chern numbers is frequently used for pairings involving a related, but different type of characteristic classes: the so-called *Chern characters*. To motivate the definition of the (total) Chern character, recall that the (total) Chern class is multiplicative under direct sums of vector bundles, i.e.

$$c(E_1 \oplus E_2) = c(E_1) \smile c(E_2). \quad (2.229)$$

Thus, the (total) Chern class translates *addition of bundles* into *multiplication in cohomology*. As a consequence, the total Chern class does not induce a natural ring homomorphism between the semi-ring $(\text{Vect}(B), \oplus, \otimes)$ of (complex) vector bundles over B (K -theory) and the ring $(H^*(B), +, \smile)$ of cohomology groups of B [15].

The (total) Chern character $[ch] \in H^*(B)$ is designed to remedy this by considering an algebraic combination of Chern classes that takes *direct sums* to *sums* and *tensor products* to *cup products*, i.e.

$$ch(E_1 \oplus E_2) = ch(E_1) + ch(E_2) \quad \text{and} \quad ch(E_1 \otimes E_2) = ch(E_1) \smile ch(E_2). \quad (2.230)$$

Specifically, the (total) Chern character may be defined as

$$ch(E) = \text{rank}(E) + \sum_{k>0} s_k(c_1(E), \dots, c_k(E))/k!, \quad (2.231)$$

where

$$s_k(c_1(E), \dots, c_k(E)) = \sum_{j=1}^k s_{k-j}(c_1(E), \dots, c_{k-j}(E)) c_j(E) \quad (2.232)$$

denotes the recursively defined Newton polynomials [15]. The first terms of Eq. (2.231) are

$$\begin{aligned} ch_1(E) &= c_1(E) \\ ch_2(E) &= \frac{1}{2} [c_1(E)^2 - 2c_2(E)] \\ ch_3(E) &= \frac{1}{6} [c_1(E)^3 - 3c_1(E)c_2(E) + 3c_3(E)] \end{aligned} \quad (2.233)$$

and called the first, second and third Chern character of E , respectively. Note that without the $1/k!$ prefactor, the k -th Chern character $ch_k(E) \propto s_k(c_1(E), \dots, c_k(E))$ is an integer linear combination of

degree- k products of Chern classes. In particular, $ch_3(E)$ contains exactly the combinations of Chern classes that we anticipated in our example Eq. (2.228) for integral “mathematical” Chern numbers. This means that, aside from the $1/k!$ prefactor, the k -th Chern character $ch_k(E)$ defines an integral cohomology class. However, the presence of the fractional $1/k!$ prefactor disrupts this notion: the full Chern characters can only take integer values if the linear combination $s_k(c_1(E), \dots, c_k(E))$ of Chern classes happens to be divisible by $k!$. This is not always the case. In fact, this is precisely the sacrifice required to make the construction of the Chern character work: one must work with rational rather than integer coefficients. Consequently, we generally have

$$[ch(E)] \in H^*(B, \mathbb{Q}) \quad \text{and} \quad [ch_k(E)] \in H^{2k}(B, \mathbb{Q}). \quad (2.234)$$

Note that $ch_1(E) \in H^2(B, \mathbb{Z})$ is always integral because $ch_1(E) = c_1(E)$ as shown in Eq. (2.233). In Chern–Weil theory, the total Chern character $[ch(\mathcal{F})] \in H^*(B, \mathbb{Q})$ is defined using the invariant polynomial

$$ch(\mathcal{F}) = \text{tr} \left(\exp \left[\frac{i\mathcal{F}}{2\pi} \right] \right) = \sum_{k=1}^{\infty} \frac{1}{k!} \text{tr} \left(\frac{i\mathcal{F}}{2\pi} \right)^k, \quad (2.235)$$

giving the expression

$$ch_k(\mathcal{F}) = \frac{1}{k!} \text{tr} \left(\frac{i\mathcal{F}}{2\pi} \right)^k \quad (2.236)$$

for the k -th Chern character $[ch_k(\mathcal{F})] \in H^{2k}(B, \mathbb{Q})$. This allows us to identify the definition

$$Ch_n(E) = \left(\frac{i}{2\pi} \right)^n \oint_B \text{tr}(\mathcal{F})^n, \quad (2.237)$$

that the authors of Ref. [14] give for the n -th Chern number $Ch_n(E)$ of a vector bundle $E \xrightarrow{\pi} B$ as the de Rham duality pairing

$$Ch_n(E) = \langle [ch_n(E)], [B] \rangle \in \mathbb{Q} \quad (2.238)$$

between the n -th Chern character $ch_n \in H^{2n}(B, \mathbb{Q})$ and the orientation class of B .²⁵ We just argued that, by definition of the Chern characters alone, this expression is only guaranteed to be rational for $n > 1$. Yet, Chern numbers are consistently advertised as integral topological invariants. One particularly beautiful reason for this is the following.

In condensed matter physics we are often concerned with principal $U(N)$ -bundles $P \xrightarrow{\pi_P} B$ associated with complex rank- N state bundles $\Psi \xrightarrow{\pi_\Psi} B$ over base spaces B that are *spin manifolds*, i.e. that admit a global lift of the structure group $SO(\dim_{\mathbb{R}}(B))$ to the Spin group $\text{Spin}(\dim_{\mathbb{R}}(B))$ on their tangent bundle $TB \xrightarrow{\pi_{TB}} B$.²⁶ Such a lift allows for the definition of a spinor bundle $S \xrightarrow{\pi_S} B$ over B . For even-dimensional base manifolds, $\dim_{\mathbb{R}}(B) = 2n$, this spinor bundle decomposes as a direct sum $S = S^+ \oplus S^-$ of left- and right-handed spinor bundles $S^+ \xrightarrow{\pi_{S^+}} B$ and $S^- \xrightarrow{\pi_{S^-}} B$, and one can define a Dirac operator

$$\not{D} : \Gamma(S^+) \rightarrow \Gamma(S^-), \quad (2.239)$$

which acts as a first-order differential operator on the sections of S . Now the principal $U(N)$ -bundle $P \xrightarrow{\pi_P} B$ from a condensed matter problem can be used to define the twisted spinor bundle $S \otimes P$, whose sections describe spinor fields interacting with an external $U(N)$ gauge degree of freedom captured by the (local) connection (form) $\mathcal{A}$ of P . The Dirac operator

$$\not{D}_{\mathcal{A}} : \Gamma(S^+ \otimes P) \rightarrow \Gamma(S^- \otimes P) \quad (2.240)$$

²⁵The authors are clearly aware of this: they call their Chern number Ch_n instead of C_n , hinting the relation to the n -th Chern character rather than the n -th Chern class.

²⁶Recall that a manifold M is spin if and only if $w_1(TM) = w_2(TM) = 0$, i.e. if the first and second Stiefel–Whitney classes of its tangent bundle vanish identically. Important examples include all spheres $\mathbb{S}^n$ and tori $\mathbb{T}^n$.

on this twisted spinor bundle is called the Dirac operator twisted by P . We write $\mathcal{D}_{\mathcal{A}}$ to highlight that $\mathcal{D}$ is twisted by P via $\mathcal{A}$. In local coordinates, $\mathcal{D}_{\mathcal{A}}$ takes the familiar form

$$\mathcal{D}_{\mathcal{A}} = \not{D} + i\not{A} = \gamma^\mu (\partial_\mu + i\mathcal{A}_\mu) , \quad (2.241)$$

where the $\gamma^\mu \in \text{Cl}(2n)$ are gamma matrices generating the Clifford algebra associated with the base manifold B . When B is compact, the twisted Dirac operator $\mathcal{D}_{\mathcal{A}}$ is elliptic [39]. In that case, one can define its analytical index as

$$\text{ind}(\mathcal{D}_{\mathcal{A}}) = \dim \ker(\mathcal{D}_{\mathcal{A}}) - \dim \ker(\mathcal{D}_{\mathcal{A}}^\dagger) , \quad (2.242)$$

and its topological index as

$$\text{top}(\mathcal{D}_{\mathcal{A}}) = \oint_B \left(\hat{A}(TB) ch(P) \right) \Big|_{\text{vol}} , \quad (2.243)$$

where $ch(P)$ denotes the (total) Chern character of P and $\hat{A}(TB)$ is a characteristic class of the tangent-bundle $TB \xrightarrow{\pi_{TB}} B$ called the $\hat{A}$ -genus or Dirac-genus of B .²⁷ The restriction $|_{\text{vol}}$ of the integrand to the volume ensures that only top-degree forms are picked up so that the integration makes sense. The famous Atiyah–Singer index theorem [47, 74] then states that

$$\text{ind}(\mathcal{D}_{\mathcal{A}}) = \text{top}(\mathcal{D}_{\mathcal{A}}) , \quad (2.244)$$

i.e. that the analytical and the topological indices of $\mathcal{D}_{\mathcal{A}}$ are the same. Suppose B has a trivial Dirac genus²⁸ of $\hat{A}(TB) = 1$, so that the topological index, Eq. (2.243), becomes

$$\text{top}(\mathcal{D}_{\mathcal{A}}) = \oint_B ch_n(P) . \quad (2.245)$$

If we additionally identify the numbers $\nu_+ = \dim \ker(\mathcal{D}_{\mathcal{A}})$ and $\nu_- = \dim \ker(\mathcal{D}_{\mathcal{A}}^\dagger)$ of positive and negative chirality zero-energy modes of $\mathcal{D}_{\mathcal{A}}$ [39] to rewrite the analytical index, Eq. (2.242), as

$$\text{ind}(\mathcal{D}_{\mathcal{A}}) = \nu_+ - \nu_- , \quad (2.246)$$

the Atiyah–Singer index theorem from Eq. (2.244) takes the simple form

$$\nu_+ - \nu_- = \oint_B ch_n(P) . \quad (2.247)$$

Now, since ν_+ and ν_- count the numbers of positive and negative chirality zero-energy modes of $\mathcal{D}_{\mathcal{A}}$, they are clearly natural numbers, i.e. $\nu_+, \nu_- \in \mathbb{N}$. The analytical index, Eq. (2.246), is therefore obviously an integer.²⁹ Consequently, Eq. (2.247) ensures that the top Chern number, Eq. (2.237), of a $U(N)$ bundle over a closed, compact spin manifold B of even dimension $\dim(B) = 2n$ is always an integer. In the special case of $U(N)=U(1)$, this quantisation corresponds to the Adler–Bell–Jackiw chiral anomaly on the even-dimensional manifold B [39, 57, 58, 75].

Variants of the Atiyah–Singer index theorem, most notably the Atiyah–Patodi–Singer index theorem [25, 48–50] for fibre bundles over base manifolds *with* boundary, provide the mathematical foundation for the much-cited bulk-boundary correspondence of topological insulators [14, 25, 47], relating topological bulk indices to the spectral flow of boundary Dirac operators [39].

²⁷Similar to how the Chern character $ch(B)$ is a rational combination of Chern classes $c_i \in H^{2i}(B, \mathbb{Z})$, the Dirac genus is a $\mathbb{Q}$ -valued polynomial $\hat{A}(TB) = 1 + -\frac{1}{24}p_1 + \frac{1}{5760}(4p_2 - 7p_1^2) + \dots$ in the Pontryagin classes $p_i \in H^{4i}(B, \mathbb{Z})$. Consequently, $\hat{A}(TB)$ can only be paired against the fundamental class $[B]$ of B to produce a characteristic number $\hat{A}(B) = \oint_B \hat{A}(TB)$ if $\dim_{\mathbb{R}}(B) = 4n$, i.e. if the dimension of B is divisible by four. Unfortunately, the characteristic number $\hat{A}(B)$ is commonly referred to as the Dirac genus as well.

²⁸The Dirac genus is trivial, $\hat{A}(TB) = 1$, if all Pontryagin classes of the tangent bundle $TB \xrightarrow{\pi_{TB}} B$ vanish, cf. first few terms of $\hat{A}(TB)$ in the previous footnote. This situation is not overly exotic. For instance, Eq. (2.223) shows that this holds if TB is flat, i.e. there exists a connection with zero curvature $\mathcal{F} = 0$ on TB . Moreover, the Pontryagin classes are stable so that $p_i(TB) = 0$ if TB is stably trivial, i.e. if there exists a trivial bundle T_{TB} such that $TB \oplus T_{TB}$ is again trivial, cf. Eq. (2.205) and discussion. The first argument, for example, shows that $\hat{A}(TB) = 1$ for all n -tori $B = \mathbb{T}^n$, while the latter ensures the same for all n -spheres $B = \mathbb{S}^n$.

²⁹The Atiyah–Singer index theorem is often regarded as the conceptual *origin* of integral topological indices of differential operators [39].

3 – Time Reversal and Particle Hole Symmetry

In quantum mechanics, a normalised quantum state $|\psi\rangle$ is an equivalence class

$$[\psi] = \{e^{i\phi} |\psi\rangle, \phi \in \mathbb{R}\} \quad (3.1)$$

of complex rays in a Hilbert space $\mathcal{H}$. The space of these equivalence classes is called the projective Hilbert space $P(\mathcal{H})$. In this framework, a symmetry transformation is an automorphism

$$S : P(\mathcal{H}) \rightarrow P(\mathcal{H}), \quad [\psi] \mapsto S[\psi], \quad (3.2)$$

that preserves the ray product $[\phi] \cdot [\psi] := |\langle\phi|\psi\rangle|$ between all states $[\phi], [\psi] \in P(\mathcal{H})$ as

$$[\phi] \cdot [\psi] = [S\psi] \cdot [S\phi]. \quad (3.3)$$

Wigner's theorem [76] states that every symmetry transformation $S : P(\mathcal{H}) \rightarrow P(\mathcal{H})$ comes either from a unitary operator $U_S : \mathcal{H} \rightarrow \mathcal{H}$ satisfying

$$\langle U_S \phi | U_S \psi \rangle = \langle \phi | U_S^\dagger U_S | \psi \rangle = \langle \phi | \psi \rangle \quad (3.4)$$

or an antiunitary operator $A_S : \mathcal{H} \rightarrow \mathcal{H}$ satisfying

$$\langle A_S \phi | A_S \psi \rangle = \langle \phi | A_S^\dagger A_S | \psi \rangle^* = \langle \phi | \psi \rangle^* = \langle \psi | \phi \rangle. \quad (3.5)$$

In second quantisation, a non-interacting, particle-number conserving quantum system is characterised by a Hamiltonian operator

$$H = h_{\alpha\beta} c_\alpha^\dagger c_\beta, \quad (3.6)$$

where c_α and $c_\alpha^\dagger$ with $\alpha = 1, \dots, d$ are fermionic annihilation and creation operators satisfying the canonical anticommutation relations

$$\{c_\alpha, c_\beta^\dagger\} = \delta_{\alpha\beta} \quad \text{and} \quad \{c_\alpha, c_\beta\} = \{c_\alpha^\dagger, c_\beta^\dagger\} = 0. \quad (3.7)$$

Note that we used an Einstein notation (implicite summation of double indices) in Eq. (3.6) to aid readability. The indices α, β are multi-indices accounting for a basis of the Hilbert space $\mathcal{H}$, e.g. $\alpha = (j, \sigma)$ where j is a site index and σ is a spin index. A generic unitary symmetry S of H is then implemented by a unitary operator U_S that transforms the elementary fermionic annihilation and creation operators as

$$U_S^\dagger c_\alpha U_S = u_{\alpha\beta} c_\beta \quad \text{and} \quad U_S^\dagger c_\alpha^\dagger U_S = u_{\alpha\beta}^* c_\beta^\dagger, \quad (3.8)$$

while preserving the anticommutator relations

$$U_S \{c_\alpha, c_\beta^\dagger\} U_S^\dagger = \{c_\alpha, c_\beta^\dagger\}, \quad (3.9)$$

and the system Hamiltonian

$$U_S H U_S^\dagger = H \quad \Longleftrightarrow \quad [H, U_S] = 0. \quad (3.10)$$

The unitary symmetries of a system form a group $\mathcal{G}_S$, whose irreducible representations can be used to bring H into a block-diagonal form.¹ In this sense, unitary symmetries allow for a systematic reduction of the problem's full dimension to that of a single irreducible block.

¹This is an application of Schur's lemma, see e.g. Ref. [77].

In Ref. [14] the authors present an exhaustive classification of these irreducible blocks in terms of three additional “non-standard” symmetries. These symmetries are time reversal symmetry (TRS) $\mathcal{T}$, particle-hole symmetry (PHS) $\mathcal{C}$, and their combination, chiral symmetry $\mathcal{S} = \mathcal{T} \cdot \mathcal{C}$. The latter has to be included because some systems are only invariant under the combined transformation $\mathcal{S}$, but not under $\mathcal{T}$ and $\mathcal{C}$ individually. Both TRS and PHS can square as

$$\mathcal{T}^2 = \pm 1 \quad \text{and} \quad \mathcal{C}^2 = \pm 1, \quad (3.11)$$

which gives rise to ten possible ways in which a given Hamiltonian H can transform under $\mathcal{T}, \mathcal{C}$ and $\mathcal{S}$. To see this, we write $\mathbf{T} = 0$ ($\mathbf{C} = 0$) if H is not invariant under TRS (PHS), and $\mathbf{T} = \pm 1$ ($\mathbf{C} = \pm 1$) if it is invariant with $\mathcal{T}^2 = \pm 1$ ($\mathcal{C}^2 = \pm 1$). Thus, there are $3 \cdot 3 = 9$ possible ways in which H can transform under combinations of $\mathcal{T}$ and $\mathcal{C}$. In all these cases, $\mathbf{S}$ is determined as $\mathbf{S} = \mathbf{T} \cdot \mathbf{C}$. The tenth possibility is then $\mathbf{T} = \mathbf{C} = 0$ but $\mathbf{S} = 1$. The tenfold way of condensed matter theory corresponds to a tenfold classification of so-called real super division algebras in mathematics – a very brief excursion is included in App. A.1 for fun. In the following, we outline some details about TRS and PHS that will be needed later.

3.1 Time Reversal Symmetry

Time reversal symmetry (TRS) is an involutive symmetry transformation, i.e. it is an automorphism $T : \mathcal{P}(\mathcal{H}) \rightarrow \mathcal{P}(\mathcal{H})$ that squares to the identity $T^2 = \mathbb{1}$. On the underlying Hilbert space $\mathcal{H}$, TRS is implemented by an antiunitary operator $\mathcal{T} : \mathcal{H} \rightarrow \mathcal{H}$, which transforms the fermionic annihilation operators as

$$\mathcal{T} c_\alpha \mathcal{T}^\dagger = U_{T\alpha\beta} c_\beta \quad \text{and} \quad \mathcal{T} i \mathcal{T}^\dagger = -i, \quad (3.12)$$

where $U_{T\alpha\beta}$ are elements of a unitary matrix U_T representing the linear part $\mathcal{U}_T$ of $\mathcal{T}$ in the c_α -basis. This means that we can write $\mathcal{T}$ as a product

$$\mathcal{T} = \mathcal{U}_T \mathcal{K} \quad (3.13)$$

of the unitary operator $\mathcal{U}_T$ and the complex conjugation operator $\mathcal{K}$ in the c_α -basis, which implements the antilinearity of $\mathcal{T}$. In the generic case of fermions on a discrete lattice with sites $j = 1, \dots, L$, TRS acts as

$$\mathcal{T} c_j \mathcal{T}^\dagger = \delta_{jk} c_k, \quad (3.14)$$

if the fermions are spinless and as

$$\mathcal{T} c_{j\gamma} \mathcal{T}^\dagger = \delta_{jk} (i\sigma_y)_{\gamma\delta} c_{k\delta}, \quad (3.15)$$

if they have spin one-half. Here, $\gamma, \delta \in \{\uparrow, \downarrow\}$ label the spin indices of the elementary annihilation and creation operators. A system with Hamiltonian H is called TRS invariant if

$$\mathcal{T} H \mathcal{T}^\dagger = H. \quad (3.16)$$

Since $T : \mathcal{P}(\mathcal{H}) \rightarrow \mathcal{P}(\mathcal{H})$ fulfils $T^2 = \mathbb{1}$, the operator $\mathcal{T}$ implementing it on $\mathcal{H}$ must satisfy $\mathcal{T}^2 = e^{i\varphi} \mathbb{1}$. Invoking the antilinearity $\mathcal{T} e^{i\alpha} = e^{-i\varphi} \mathcal{T}$ of $\mathcal{T}$ together with

$$\mathcal{T} e^{i\varphi} = \mathcal{T} e^{i\varphi} \mathbb{1} = \mathcal{T} \mathcal{T}^2 = \mathcal{T}^2 \mathcal{T} = e^{i\varphi} \mathbb{1} \mathcal{T} = e^{i\varphi} \mathcal{T} \quad (3.17)$$

yields the constraint $e^{i\varphi} = e^{-i\varphi}$, i.e. $e^{i\varphi} = \pm 1$. We can express this constraint in terms of the unitary matrix representation U_T of $\mathcal{U}_T$ by substituting Eq. (3.6) into Eq. (3.16). This yields

$$U_T^\dagger h^* U_T = h. \quad (3.18)$$

If we take the complex conjugate $U_T^\dagger h U_T^* = h^*$ of this and plug it back into Eq. (3.18), we obtain

$$(U_T^* U_T)^\dagger h (U_T^* U_T) = h. \quad (3.19)$$

Thus, $U_T^* U_T$ commutes with the Hamiltonian matrix h . Since we are working within an irreducible representation of the symmetry class of all TRS-invariant Hamiltonians, Schur's lemma implies that [14]

$$U_T^* U_T = e^{i\varphi} \mathbb{1}. \quad (3.20)$$

Using the unitarity of U_T , we can multiply this by $U_T^\dagger$ from the left and by $U_T^\dagger$ from the right to get

$$\mathbb{1} = e^{i\varphi} U_T^\dagger U_T^\dagger \implies U_T^\dagger U_T^\dagger = e^{-i\varphi} \mathbb{1}, \quad (3.21)$$

the transpose of which is

$$U_T^* U_T = e^{-i\varphi} \mathbb{1}. \quad (3.22)$$

Thus, we find

$$e^{i\varphi} = e^{-i\varphi} \implies e^{i\varphi} = \pm 1, \quad (3.23)$$

and, accordingly,

$$U_T^* U_T = \pm \mathbb{1}. \quad (3.24)$$

As a result, we get

$$\mathcal{T}^2 c_\alpha \mathcal{T}^{\dagger 2} = \mathcal{T} U_{T\alpha\beta} c_\beta \mathcal{T}^\dagger = U_{T\alpha\beta}^* U_{T\beta\gamma} c_\gamma = \pm \delta_{\alpha\gamma} c_\gamma = \pm c_\alpha \quad (3.25)$$

and, more generally,

$$\mathcal{T}^2 O \mathcal{T}^{\dagger 2} = (\pm 1)^n O \quad (3.26)$$

for an operator O consisting of n fermionic annihilation and creation operators. Consider a non-interacting particle-number conserving Hamiltonian as in Eq. (3.6). The Fock space $\mathcal{F}(\mathcal{H})$ of $\mathcal{H}$ is formed by the action of polynomials in the single-fermion creation operators $c_\alpha^\dagger$ on the vacuum state $|0\rangle$, which is defined via $c_\alpha |0\rangle = 0$ for all $\alpha = 1, \dots, d$. A generic Fock state in $\mathcal{F}(\mathcal{H})$ takes the form

$$|n_1, \dots, n_d\rangle = \prod_{\alpha=1}^d (c_\alpha^\dagger)^{n_\alpha} |0\rangle. \quad (3.27)$$

With $\mathcal{T} |0\rangle = |0\rangle$ we then find

$$\mathcal{T}^2 |n_1, \dots, n_d\rangle = \mathcal{T}^2 \left[\prod_{\alpha=1}^d (c_\alpha^\dagger)^{n_\alpha} \right] \mathcal{T}^{\dagger 2} \mathcal{T}^2 |0\rangle = (\pm 1)^N |0\rangle, \quad (3.28)$$

where $N = \sum_{\alpha=1}^d n_\alpha$ is the total number of fermions in $|n_1, \dots, n_d\rangle$. We can therefore use the total fermion number operator $\mathcal{N} = \sum_{\alpha=1}^d c_\alpha^\dagger c_\alpha$ to write $\mathcal{T}^2$ as

$$\mathcal{T}^2 = (\pm 1)^\mathcal{N}. \quad (3.29)$$

An important consequence in systems with $U_T^* U_T = -\mathbb{1}$ and $\mathcal{T}^2 = (-1)^\mathcal{N}$ is that TRS invariance $[\mathcal{T}, H] = 0$ leads to a degeneracy of the spectrum. To see this, we consider a single-particle energy eigenstate $|n\rangle$ of a TRS invariant Hamiltonian H ,

$$H |n\rangle = E_n |n\rangle, \quad (3.30)$$

and note that its time-reversed partner $|\mathcal{T}n\rangle \equiv \mathcal{T} |n\rangle$ is an eigenstate with the same energy

$$H |\mathcal{T}n\rangle = H \mathcal{T} |n\rangle = \mathcal{T} H |n\rangle = \mathcal{T} E_n |n\rangle = E_n \mathcal{T} |n\rangle = E_n |\mathcal{T}n\rangle. \quad (3.31)$$

The pair $\{|n\rangle, |\mathcal{T}n\rangle\}$ is called a Kramers pair. A celebrated theorem by Kramers states that the two states of a Kramers pair are orthogonal if $\mathcal{T}^2 = -\mathbb{1}$. This follows from

$$\mathcal{T}^2 = -\mathbb{1} \implies \mathcal{T} = -\mathcal{T}^\dagger, \quad (3.32)$$

with which

$$\langle n|\mathcal{T}n\rangle = \langle n|\mathcal{T}|n\rangle \stackrel{(\diamond)}{=} \langle n|\mathcal{T}^\dagger|n\rangle = -\langle n|\mathcal{T}|n\rangle = -\langle n|\mathcal{T}n\rangle \implies \langle n|\mathcal{T}n\rangle = 0. \quad (3.33)$$

In $(\diamond)$ we used that $\mathcal{T}$ is an antilinear operator, so it satisfies $\langle\psi|\mathcal{T}|\phi\rangle = \langle\phi|\mathcal{T}^\dagger|\psi\rangle$. An important class of systems with $\mathcal{T}^2 = (-\mathbb{1})^\mathcal{N}$ are spin one-half fermions on a discrete lattice. We can show this directly by using $U_{T(j\gamma)(k\xi)} = \delta_{jk}(i\sigma_y)_{\gamma\xi}$ from Eq. (3.15) to determine the matrix elements of $U_T^*U_T$ as

$$\begin{aligned} \left(U_T^*U_T\right)_{(j\gamma)(k\xi)} &= U_{T(j\gamma)(f\eta)}^* U_{(f\eta)(k\xi)} \\ &= \delta_{jf}(i\sigma_y)_{\gamma\eta} \delta_{fk}(i\sigma_y)_{\eta\xi} \\ &= [\delta_{jf}\delta_{fk}] [(i\sigma_y)_{\gamma\eta}(i\sigma_y)_{\eta\xi}] \\ &= \delta_{jk} [-\delta_{\gamma\xi}] \\ &= -\delta_{(j\gamma)(k\xi)} \\ &= (-\mathbb{1})_{(j\gamma)(k\xi)}, \end{aligned} \quad (3.34)$$

which readily confirms $U_T^*U_T = -\mathbb{1}$ and hence $\mathcal{T}^2 = (-\mathbb{1})^\mathcal{N}$ for these systems. Similarly, we note that we trivially have $U_T^*U_T = \mathbb{1}$ in systems of spinless fermions on a discrete lattice, cf. Eq. (3.14).

3.2 Particle Hole Symmetry

Particle hole symmetry (PHS) is an involutive symmetry transformation $C : \mathcal{P}(\mathcal{H}) \rightarrow \mathcal{P}(\mathcal{H})$ that satisfies $C^2 = \mathbb{1}$. On the underlying Hilbert space $\mathcal{H}$, PHS is implemented by an antiunitary operator $\mathcal{C} : \mathcal{H} \rightarrow \mathcal{H}$ that transforms the elementary fermionic annihilation operators as

$$\mathcal{C}c_\alpha\mathcal{C}^\dagger = U_{C\alpha\beta}c_\beta^\dagger \quad \text{and} \quad \mathcal{C}i\mathcal{C}^\dagger = -i, \quad (3.35)$$

where $U_{C\alpha\beta}$ are elements of a matrix U_C , which represents a projective involution operator $\mathcal{U}_C : \mathcal{H} \rightarrow \mathcal{H}$ with $\mathcal{U}_C^2 = e^{i\varphi}\mathbb{1}$ in the c_α -basis. The central aspect of the PHS transformation is that annihilation operators are mapped to creation operators and vice versa. We can implement this via an antiunitary particle hole conjugation (PHC) operation, which is defined via

$$\bar{\Xi}c_\alpha\bar{\Xi}^\dagger = \delta_{\alpha\beta}c_\beta^\dagger \quad \text{and} \quad \bar{\Xi}i\bar{\Xi}^\dagger = -i. \quad (3.36)$$

PHC is also called charge conjugation because it flips the sign, $\bar{\Xi}\mathcal{Q}\bar{\Xi}^\dagger = -\mathcal{Q}$, of the $U(1)$ charge operator $\mathcal{Q} = \mathcal{N} - d/2$ where $\mathcal{N} = \sum_{\alpha=1}^d c_\alpha^\dagger c_\alpha$ is the fermion number operator and $d/2$ is half the single-particle Hilbert space dimension. The PHC operator $\bar{\Xi}$ is antiunitary because it comes from the antilinear Fréchet-Riesz isomorphism $\mathcal{H} \ni |\psi\rangle \mapsto \langle\psi| \in \mathcal{H}^*$ on the single particle Hilbert space $\mathcal{H}$. With $\bar{\Xi}$, we can write $\mathcal{C}$ as a product

$$\mathcal{C} = \mathcal{U}_C\bar{\Xi} = \mathcal{U}_C\bar{\Xi}\mathcal{K} \quad (3.37)$$

of the involution operator $\mathcal{U}_C$ and the antiunitary PHC operator $\bar{\Xi}$, or, alternatively, the involution operator $\mathcal{U}_C$, the unitary part of the PHC operator $\bar{\Xi}$ and the complex conjugation operator $\mathcal{K}$ in the c_α -basis. For trivial $\mathcal{U}_C = \mathbb{1}$, the PHS operator $\mathcal{C}$ is simply given by the antiunitary PHC operator $\bar{\Xi}$. However, $\mathcal{C} = \bar{\Xi}$ can never be a meaningful symmetry of non-interacting quadratic Hamiltonians² like the one given in Eq. (3.6). In fact, any such Hamiltonian is odd under the antiunitary PHC transformation

$$\bar{\Xi}H\bar{\Xi}^\dagger \stackrel{(*)}{=} -H, \quad (3.38)$$

²This includes all BdG Hamiltonians, cf. Sec. 5.

so H and $\bar{\Xi}$ anticommute as $\{H, \bar{\Xi}\} = 0$. Moreover, this transformation behaviour is a direct result of the canonical anticommutation relations of the $c_\alpha^\dagger$ and c_β . For this reason, some authors refer to Eq. (3.38) as a tautological PHC symmetry property [78]. A meaningful PHS invariance

$$\mathcal{C}H\mathcal{C}^\dagger = H \quad (3.39)$$

therefore requires a non-trivial unitary operator $\mathcal{U}_C$. Indeed, we find that

$$\mathcal{C}H\mathcal{C}^\dagger = \mathcal{U}_C \bar{\Xi} H \bar{\Xi}^\dagger \mathcal{U}_C^\dagger = -\mathcal{U}_C H \mathcal{U}_C^\dagger \stackrel{!}{=} H \implies \{H, \mathcal{U}_C\} \stackrel{!}{=} 0, \quad (3.40)$$

i.e. that $\mathcal{U}_C$ has to undo the sign reversal due to the bare particle hole conjugation. It follows that, $\mathcal{U}_C$ swaps the positive and negative energy eigenstates of H . In condensed matter, $\mathcal{U}_C$ typically comes from a (sub)lattice transformation in real space or a momentum shift in k -space. A simple example is the tight-binding Hamiltonian

$$H = t \sum_{j=1}^{L-1} \left(c_j^\dagger c_{j+1} + c_{j+1}^\dagger c_j \right) \quad (3.41)$$

of spinless fermions c_j on a one-dimensional chain of L sites, which is invariant under a PHS transformation

$$\mathcal{C}c_j\mathcal{C}^\dagger = (-1)^j c_j^\dagger, \quad (3.42)$$

with the involution matrix $U_{Cjk} = (-1)^j \delta_{jk}$. Just like with TRS, the involutive property $\mathcal{C}^2 = \mathbb{1}$ of $\mathcal{C} : \mathcal{P}(\mathcal{H}) \rightarrow \mathcal{P}(\mathcal{H})$ means that $\mathcal{C} : \mathcal{H} \rightarrow \mathcal{H}$ squares to $\mathcal{C}^2 = e^{i\varphi} \mathbb{1}$ for some phase factor $e^{i\varphi}$ and again the antiunitarity of $\mathcal{C}$ restricts the possible phase factors to $e^{i\varphi} = \pm 1$. Thus, we have

$$\mathcal{C}^2 = \pm \mathbb{1}. \quad (3.43)$$

The particle-hole operator from Eq. (3.42) satisfies $\mathcal{C}^2 = +1$.

4 – Geometric Phases

Whenever a physical system depends on parameters λ from a parameter manifold Λ , a cyclic evolution of these parameters may induce a phase shift that depends solely on the path traversed in Λ . This phenomenon is known as a geometric phase. Mathematically, geometric phases are closely related to the concept of holonomy in fibre bundles, which we addressed in Sec. 2.2.5. The fibre-bundle perspective enables the unification of a wide range of physical phenomena and provides a natural link between physics, differential geometry, and topology.

In this section, we discuss the fundamental concepts of Abelian and non-Abelian geometric phases in physics, with a particular emphasis on quantum mechanical applications. Throughout, we will draw on the terminology of fiber bundle theory developed in Sec. 2.2. We begin with a brief excursion into classical physics and review the Foucault pendulum as an introductory example. After that, we give a quantum mechanical derivation of Berry’s adiabatic phase and identify it as an Abelian $U(1)$ holonomy. As a natural progression, we then extend the notion of an Abelian geometric $U(1)$ phase to that of a non-Abelian geometric $U(N)$ phase known as Wilczek–Zee phase. Finally, we discuss a generalised perspective on geometric phases that is due to Aharonov and Anandan [79]. Throughout this chapter we closely follow Refs. [80] and [39].

4.1 The Foucault Pendulum

The Foucault pendulum is an experiment that was devised to demonstrate the rotation of earth. Foucault observed that a rotating earth would cause the plane in which a pendulum swings to precess over time. At a given latitude α , it would only return to its original orientation after $T = 24/|\sin(\alpha)|$ hours. This is most obvious at the poles ($\alpha = \pm\pi/2$), where the pendulum, and hence its plane of oscillation, trivially rotates around its own suspension axis by 2π every 24 hours. Note that the formula for T diverges at the equator, where $\alpha = 0$ and $T \rightarrow \infty$. This reflects the fact that the precession rate (angular speed), $|\omega| = 2\pi/T$, of the Foucault pendulum vanishes at the equator, meaning that the swing plane there remains fixed and does not precess at all.

To understand this effect in terms of parallel transport, we adopt an alternative but equivalent viewpoint. Instead of a rotating earth with a stationary pendulum, we imagine a stationary earth on which the pendulum is transported around the latitude circle $\ell(\alpha)$ at constant speed, completing one full circle after 24 hours. For conceptual clarity, we represent earth as a two-sphere $\mathbb{S}^2$ and assume that the pendulum is infinitely long, so that the plane in which its tip moves is locally flat and always tangent to the earth’s surface at the point of suspension. Each swing direction can then be regarded as a tangent vector to a two-sphere $\mathbb{S}^2$ representing earth. By consistently orienting these swing vectors we obtain a vector field X describing the pendulum’s swing directions at each point. As the pendulum moves along $\ell(\alpha)$, its swing direction evolves according to parallel transport: the swing vector X is moved along the latitude circle without twisting relative to the sphere’s surface. Figure 4.1 illustrates how the parallel transport of the swing direction vector $X(t)$ along a latitude of $\mathbb{S}^2$ can change its orientation. Applying this analysis to the Foucault pendulum yields the anticipated phase shift of $2\pi \sin(\alpha)$ after one complete circuit around the earth.

The geometric precession of the Foucault pendulum only becomes dominant if there is an adiabatic separation of time scales. It is often said, that the rotation speed must be slow compared to the pendulum’s oscillation time scale. This may sound a bit curious. After all, the earth is rotating very fast – hundreds of meters per second at most latitudes. So how is the Foucault setup adiabatic enough to be observable? The answer is that the relevant time scale separation is between the pendulum’s *rapid oscillation speed*, completing a swing in a few seconds, and the *much slower precession of its swing direction due to earth’s rotation*, which happens over the course of many hours.¹

¹Recall that the shortest possible period T of the swing direction is realised at the poles where $T = 24$ hours.

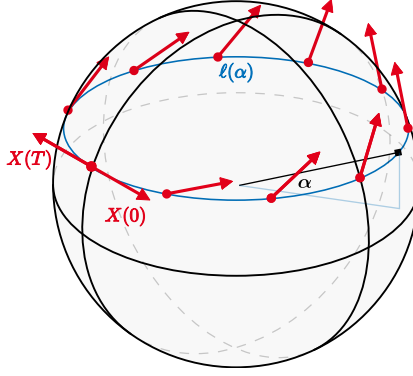


Figure 4.1: Parallel transport of the Foucault pendulum's swing direction X along the latitude circle $\ell(\alpha)$ for $\alpha = \pi/6 = 30^\circ$.

4.2 The Berry Phase

The state of a quantum mechanical system is represented by a ray in a Hilbert space $\mathcal{H}$, i.e. it corresponds to an equivalence class

$$[|\psi\rangle] \equiv \{|\phi\rangle = c |\psi\rangle, c \in \mathbb{C}^*\} \quad (4.1)$$

of all states $|\phi\rangle$ that differ from $|\psi\rangle$ only by a non-zero complex prefactor $c \in \mathbb{C}^*$. Of course, we usually work with normalised quantum states, so that the equivalence class is reduced to

$$[|\psi\rangle] \equiv \{|\phi\rangle = e^{i\phi} |\psi\rangle, \phi \in \mathbb{R}\}, \quad (4.2)$$

and we say that quantum states are only defined up to a global phase. Since that phase cancels out in expectation values, it is generally considered unphysical. However, Berry demonstrated in Ref. [81] that this phase may have observable consequences when a quantum system undergoes a cyclic adiabatic evolution. For the next part we closely follow Ref. [39].

Consider a quantum system with a Hamiltonian $H(\boldsymbol{\lambda})$ that depends on a set of parameters $\boldsymbol{\lambda}$ from some parameter manifold Λ . If $\boldsymbol{\lambda} = \boldsymbol{\lambda}(t)$ change as a function of time, the time evolution of an isolated quantum state $|\psi(t)\rangle$ is determined by the Schrödinger equation

$$i \frac{d}{dt} |\psi(t)\rangle = H(\boldsymbol{\lambda}(t)) |\psi(t)\rangle. \quad (4.3)$$

Here, a quantum state is called isolated if it is non-degenerate for all $\boldsymbol{\lambda} \in \Lambda$, or at least for all $\boldsymbol{\lambda}(t)$ along the trajectory in Λ . If we assume that the parameters $\boldsymbol{\lambda}(t)$ change slowly enough for the adiabatic theorem to apply, then a system initially in the n -th eigenstate,

$$|\psi(0)\rangle = |n(\boldsymbol{\lambda}(0))\rangle, \quad (4.4)$$

remains in the instantaneous n -th eigenstate at all later times. Let us take Eq. (4.4) as an initial state and determine its time evolution under Eq. (4.3). The implicit time-dependence of the instantaneous eigenstates $|n(\boldsymbol{\lambda}(t))\rangle$ means that the naive ansatz

$$|\psi(t)\rangle = \exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau \right] |n(\boldsymbol{\lambda}(t))\rangle \quad (4.5)$$

does not work because

$$\begin{aligned} i \frac{d}{dt} |\psi(t)\rangle &= i \frac{d}{dt} \exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau \right] |n(\boldsymbol{\lambda}(t))\rangle \\ &= E_n(\boldsymbol{\lambda}(t)) \exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau \right] |n(\boldsymbol{\lambda}(t))\rangle + i \exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau \right] \frac{d}{dt} |n(\boldsymbol{\lambda}(t))\rangle \\ &= E_n(\boldsymbol{\lambda}(t)) |\psi(t)\rangle + \exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau \right] i \frac{d}{dt} |n(\boldsymbol{\lambda}(t))\rangle, \end{aligned} \quad (4.6)$$

whereas

$$H(\boldsymbol{\lambda}(t)) |\psi(t)\rangle = E_n(\boldsymbol{\lambda}(t)) |\psi(t)\rangle, \quad (4.7)$$

so the left-hand side of Eq. (4.3) has an extra term

$$\exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau \right] i \frac{d}{dt} |n(\boldsymbol{\lambda}(t))\rangle. \quad (4.8)$$

To compensate for this term, we introduce an additional time-dependent phase $\eta(t)$ to our ansatz, so that

$$|\psi(t)\rangle = \exp \left[-i \left(\int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau - \eta(t) \right) \right] |n(\boldsymbol{\lambda}(t))\rangle. \quad (4.9)$$

The updated $|\psi(t)\rangle$ still satisfies Eq. (4.7), but the left-hand side of Eq. (4.3) changes to

$$\begin{aligned} i \frac{d}{dt} |\psi(t)\rangle &= i \frac{d}{dt} \exp \left[-i \left(\int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau - \eta(t) \right) \right] |n(\boldsymbol{\lambda}(t))\rangle \\ &= \left(E_n(\boldsymbol{\lambda}(t)) - \frac{d\eta(t)}{dt} \right) |\psi(t)\rangle + i \exp \left[-i \left(\int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau - \eta(t) \right) \right] \frac{d}{dt} |n(\boldsymbol{\lambda}(t))\rangle. \end{aligned} \quad (4.10)$$

For Eq. (4.9) to satisfy the time-dependent Schrödinger equation, we must choose $\eta(t)$ such that

$$E_n(\boldsymbol{\lambda}(t)) |\psi(t)\rangle \stackrel{!}{=} \left(E_n(\boldsymbol{\lambda}(t)) - \frac{d\eta(t)}{dt} \right) |\psi(t)\rangle + i \exp \left[-i \left(\int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau - \eta(t) \right) \right] \frac{d}{dt} |n(\boldsymbol{\lambda}(t))\rangle, \quad (4.11)$$

which, upon cancelling $E_n(\boldsymbol{\lambda}(t)) |\psi(t)\rangle$, becomes

$$\frac{d\eta(t)}{dt} |\psi(t)\rangle \stackrel{!}{=} i \exp \left[-i \left(\int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau - \eta(t) \right) \right] \frac{d}{dt} |n(\boldsymbol{\lambda}(t))\rangle. \quad (4.12)$$

We can multiply this from the left by $\langle\psi(t)|$ as given in Eq. (4.9) to get

$$\frac{d\eta(t)}{dt} \stackrel{!}{=} i \langle n(\boldsymbol{\lambda}(t)) | \frac{d}{dt} |n(\boldsymbol{\lambda}(t))\rangle, \quad (4.13)$$

so that integration yields

$$\begin{aligned} \eta(t) &\stackrel{!}{=} i \int_0^t \langle n(\boldsymbol{\lambda}(\tau)) | \frac{d}{d\tau} |n(\boldsymbol{\lambda}(\tau))\rangle d\tau \\ &\stackrel{(\diamond)}{=} i \int_0^t \langle n(\boldsymbol{\lambda}(\tau)) | \left(\sum_j \frac{\partial}{\partial \lambda_j} |n(\boldsymbol{\lambda}(\tau))\rangle \frac{d\lambda_j(\tau)}{d\tau} \right) d\tau \\ &= i \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \sum_j \langle n(\boldsymbol{\lambda}) | \frac{\partial}{\partial \lambda_j} |n(\boldsymbol{\lambda})\rangle d\lambda_j \\ &\stackrel{(\star)}{=} i \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \langle n(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} |n(\boldsymbol{\lambda})\rangle d\boldsymbol{\lambda} \\ &\stackrel{(*)}{=} i \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \langle n(\boldsymbol{\lambda}) | d |n(\boldsymbol{\lambda})\rangle. \end{aligned} \quad (4.14)$$

Here, we have spelled out the derivative in $(\diamond)$ to show that $\eta(t)$ is independent of time as long as the adiabatic theorem applies. In $(\star)$ we introduced the notation $\partial_{\boldsymbol{\lambda}}$ for the gradient in the parameter manifold Λ , and in $(*)$ we identified the total derivative

$$d |n(\boldsymbol{\lambda})\rangle = \partial_{\boldsymbol{\lambda}} |n(\boldsymbol{\lambda})\rangle d\boldsymbol{\lambda}. \quad (4.15)$$

Note that $\eta(t)$ is real so it defines a proper phase. This can be seen using

$$\langle n(\boldsymbol{\lambda}(\tau)) | n(\boldsymbol{\lambda}(\tau)) \rangle = 1 \implies \frac{d}{d\tau} \langle n(\boldsymbol{\lambda}(\tau)) | n(\boldsymbol{\lambda}(\tau)) \rangle = 0, \quad (4.16)$$

with which

$$\begin{aligned}
0 &= \frac{d}{d\tau} \langle n(\boldsymbol{\lambda}(\tau)) | n(\boldsymbol{\lambda}(\tau)) \rangle \\
&= \left(\frac{d}{d\tau} \langle n(\boldsymbol{\lambda}(\tau)) | \right) | n(\boldsymbol{\lambda}(\tau)) \rangle + \langle n(\boldsymbol{\lambda}(\tau)) | \frac{d}{d\tau} | n(\boldsymbol{\lambda}(\tau)) \rangle \\
&= \left[\langle n(\boldsymbol{\lambda}(\tau)) | \left(\frac{d}{d\tau} | n(\boldsymbol{\lambda}(\tau)) \rangle \right) \right]^* + \langle n(\boldsymbol{\lambda}(\tau)) | \frac{d}{d\tau} | n(\boldsymbol{\lambda}(\tau)) \rangle \\
&= 2\text{Re} \left(\langle n(\boldsymbol{\lambda}(\tau)) | \frac{d}{d\tau} | n(\boldsymbol{\lambda}(\tau)) \rangle \right), \tag{4.17}
\end{aligned}$$

such that the integrand of Eq. (4.17) satisfies

$$\text{Im}(i \langle n(\boldsymbol{\lambda}(\tau)) | \frac{d}{d\tau} | n(\boldsymbol{\lambda}(\tau)) \rangle) = \text{Re}(i \langle n(\boldsymbol{\lambda}(\tau)) | \frac{d}{d\tau} | n(\boldsymbol{\lambda}(\tau)) \rangle) = 0. \tag{4.18}$$

To see whether $\eta(t)$ can be observable, we test its transformation behaviour under gauge transformations

$$|n(\boldsymbol{\lambda})\rangle \mapsto |n'(\boldsymbol{\lambda})\rangle := |n(\boldsymbol{\lambda})\rangle e^{i\chi(\boldsymbol{\lambda})}, \tag{4.19}$$

where $\chi(\boldsymbol{\lambda})$ is a smooth function of $\boldsymbol{\lambda}$. As a preparation, we determine the transformation behaviour of the integrand

$$\begin{aligned}
\langle n'(\boldsymbol{\lambda}) | d | n'(\boldsymbol{\lambda}) \rangle &= \langle n'(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n'(\boldsymbol{\lambda}) \rangle d\boldsymbol{\lambda} \\
&= e^{-i\chi(\boldsymbol{\lambda})} \langle n(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n(\boldsymbol{\lambda}) \rangle e^{i\chi(\boldsymbol{\lambda})} d\boldsymbol{\lambda} \\
&= e^{-i\chi(\boldsymbol{\lambda})} \langle n(\boldsymbol{\lambda}) | \left[(\partial_{\boldsymbol{\lambda}} | n(\boldsymbol{\lambda}) \rangle) e^{i\chi(\boldsymbol{\lambda})} + | n(\boldsymbol{\lambda}) \rangle (i \partial_{\boldsymbol{\lambda}} \chi(\boldsymbol{\lambda})) e^{i\chi(\boldsymbol{\lambda})} \right] d\boldsymbol{\lambda} \\
&= \left[\langle n(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n(\boldsymbol{\lambda}) \rangle + i (\partial_{\boldsymbol{\lambda}} \chi(\boldsymbol{\lambda})) \langle n(\boldsymbol{\lambda}) | n(\boldsymbol{\lambda}) \rangle \right] d\boldsymbol{\lambda} \\
&= \left[\langle n(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n(\boldsymbol{\lambda}) \rangle + i \partial_{\boldsymbol{\lambda}} \chi(\boldsymbol{\lambda}) \right] d\boldsymbol{\lambda}, \tag{4.20}
\end{aligned}$$

which yields

$$\begin{aligned}
\eta'(t) &= i \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \langle n'(\boldsymbol{\lambda}) | d | n'(\boldsymbol{\lambda}) \rangle \\
&= i \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \left[\langle n(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n(\boldsymbol{\lambda}) \rangle + i \partial_{\boldsymbol{\lambda}} \chi(\boldsymbol{\lambda}) \right] d\boldsymbol{\lambda} \\
&= i \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \langle n(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n(\boldsymbol{\lambda}) \rangle d\boldsymbol{\lambda} - \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} (\partial_{\boldsymbol{\lambda}} \chi(\boldsymbol{\lambda})) d\boldsymbol{\lambda} \\
&= \eta(t) - [\chi(\boldsymbol{\lambda}(t)) - \chi(\boldsymbol{\lambda}(0))], \tag{4.21}
\end{aligned}$$

where we used that $d\chi = \partial_{\boldsymbol{\lambda}} \chi d\boldsymbol{\lambda}$ is a total derivative. Now, suppose the system evolves slowly for a time T . When $\boldsymbol{\lambda}(0) \neq \boldsymbol{\lambda}(T)$, i.e. when the final parameter configuration is different from the initial one, Eq. (4.21) tells us that the additional phase $\eta(T)$ can be removed by a gauge transformation with

$$\chi(\boldsymbol{\lambda}(T)) - \chi(\boldsymbol{\lambda}(0)) = \eta(T). \tag{4.22}$$

However, this is not possible for $\boldsymbol{\lambda}(0) = \boldsymbol{\lambda}(T)$, because in that case we get

$$\chi(\boldsymbol{\lambda}(T)) - \chi(\boldsymbol{\lambda}(0)) = \chi(\boldsymbol{\lambda}(0)) - \chi(\boldsymbol{\lambda}(0)) = 0, \tag{4.23}$$

and the phase $\eta(T)$ becomes observable. Moreover, $\eta(T)$ depends only on the closed trajectory $C \subset \Lambda$, so we write $\eta = \eta(C)$. The geometric phase factor

$$u(C) \equiv \exp[i\eta(C)] = \exp \left[- \oint_C \langle n(\boldsymbol{\lambda}) | d | n(\boldsymbol{\lambda}) \rangle \right], \tag{4.24}$$

is called a Berry phase factor. The phase $\eta(C)$ is accordingly known as Berry phase. But why would we expect a non-trivial Berry phase? And even if we find a non-trivial Berry phase – what does it signify?

The answer to these questions lies in the inherent geometrical structure of the quantum system: the Berry phase reflects the way in which the quantum states are organised over the parameter manifold Λ . The natural framework to describe such structures is that of fibre bundles. In the language of fibre bundles, the equivalence classes of parameter-dependent normalised states

$$[|n(\boldsymbol{\lambda})\rangle] = \{u |n(\boldsymbol{\lambda})\rangle \mid u \in \text{U}(1)\} \quad (4.25)$$

define a principal $\text{U}(1)$ -bundle $P \xrightarrow{\pi} \Lambda$ over the parameter manifold Λ . This bundle captures the $\text{U}(1)$ degree of freedom of the states in Eq. (4.25). The projection map of P is naturally defined as

$$\pi(u |n(\boldsymbol{\lambda})\rangle) := \boldsymbol{\lambda}, \quad (4.26)$$

and the canonical local trivialisation of P is

$$\phi : V \times \text{U}(1) \rightarrow \pi^{-1}(V), (\boldsymbol{\lambda}, u |n(\boldsymbol{\lambda})\rangle) \mapsto u |n(\boldsymbol{\lambda})\rangle. \quad (4.27)$$

Fixing the phase of $|n(\boldsymbol{\lambda})\rangle$ at each point in Λ amounts to choosing a global² section

$$\sigma : \Lambda \rightarrow P, \quad \boldsymbol{\lambda} \mapsto \sigma(\boldsymbol{\lambda}) = u(\boldsymbol{\lambda}) |n(\boldsymbol{\lambda})\rangle, \quad (4.28)$$

as specified in Tab. 2.2. Note that we often write

$$|n(\boldsymbol{\lambda})\rangle \equiv u(\boldsymbol{\lambda}) |n(\boldsymbol{\lambda})\rangle, \quad (4.29)$$

because we usually work with a fixed gauge anyway. Furthermore, global section above takes the familiar form

$$u(\boldsymbol{\lambda}) |n(\boldsymbol{\lambda})\rangle = e^{i\phi(\boldsymbol{\lambda})} |n(\boldsymbol{\lambda})\rangle, \quad (4.30)$$

if we explicitly spell out $u(\boldsymbol{\lambda}) \in \text{U}(1)$ as a phase factor. Moving forward we follow the above convention and write

$$|n(\boldsymbol{\lambda})\rangle \equiv \sigma(\boldsymbol{\lambda}) = u(\boldsymbol{\lambda}) |n(\boldsymbol{\lambda})\rangle \quad (4.31)$$

for better readability. Having identified the bundle structure of P , we may equip it with a connection. Specifically, we define the Berry connection

$$\mathcal{A} := \langle n(\boldsymbol{\lambda}) | d |n(\boldsymbol{\lambda})\rangle =: \mathcal{A}_\mu(\boldsymbol{\lambda}) d\lambda_\mu, \quad (4.32)$$

where $d = (\partial/\partial\lambda_\mu) d\lambda_\mu$ denotes the exterior derivative in Λ . To show that Eq. (4.32) defines an Ehresmann connection on P we first note that $\mathcal{A}$ is skew-Hermitian,

$$0 = d \langle n(\boldsymbol{\lambda}) | n(\boldsymbol{\lambda}) \rangle = (d \langle n(\boldsymbol{\lambda}) |) |n(\boldsymbol{\lambda})\rangle + \langle n(\boldsymbol{\lambda}) | d |n(\boldsymbol{\lambda})\rangle = \langle n(\boldsymbol{\lambda}) | d |n(\boldsymbol{\lambda})\rangle^* + \langle n(\boldsymbol{\lambda}) | d |n(\boldsymbol{\lambda})\rangle, \quad (4.33)$$

and therefore an element of the Lie algebra $\mathfrak{u}(1)$ of $\text{U}(1)$. This makes $\mathcal{A}$ a $\mathfrak{u}(1)$ -valued one-form on Λ , cf. Def. 2.2.9. Next, we must show that $\mathcal{A}$ satisfies the compatibility condition given in Eq. (2.154). To this end, we consider two overlapping charts $U_j, U_k \subset \Lambda$ with local sections

$$|n(\boldsymbol{\lambda})\rangle_j \equiv \sigma_j(\boldsymbol{\lambda}) \quad \text{and} \quad |n(\boldsymbol{\lambda})\rangle_k \equiv \sigma_k(\boldsymbol{\lambda}), \quad (4.34)$$

that are related by

$$|n(\boldsymbol{\lambda})\rangle_k = |n(\boldsymbol{\lambda})\rangle_j t_{jk}(\boldsymbol{\lambda}), \quad (4.35)$$

where $t_{jk}(\boldsymbol{\lambda}) \in \text{U}(1)$ are the transition functions between the two sections σ_j and σ_k . With this, we find

$$\begin{aligned} \mathcal{A}_k &= {}_k\langle n(\boldsymbol{\lambda}) | d |n(\boldsymbol{\lambda})\rangle_k \\ &= t_{jk}(\boldsymbol{\lambda})^{-1} {}_j\langle n(\boldsymbol{\lambda}) | d |n(\boldsymbol{\lambda})\rangle_j t_{jk}(\boldsymbol{\lambda}) \\ &= t_{jk}(\boldsymbol{\lambda})^{-1} {}_j\langle n(\boldsymbol{\lambda}) | \left[d |n(\boldsymbol{\lambda})\rangle_j t_{jk}(\boldsymbol{\lambda}) + |n(\boldsymbol{\lambda})\rangle_j dt_{jk}(\boldsymbol{\lambda}) \right] \\ &= t_{jk}(\boldsymbol{\lambda})^{-1} \mathcal{A}_j t_{jk}(\boldsymbol{\lambda}) + t_{jk}(\boldsymbol{\lambda})^{-1} dt_{jk}(\boldsymbol{\lambda}), \end{aligned} \quad (4.36)$$

²Note that this bundle need not admit a global section, in that case we have to work in local sections and make use of transition functions as explained in Sec. (2.2).

which readily reproduces Eq. (2.154). Moreover, we may spell out the $t_{jk}(\boldsymbol{\lambda}) \in \text{U}(1)$ as $t_{jk}(\boldsymbol{\lambda}) = e^{i\chi(\boldsymbol{\lambda})}$ and utilise the fact that $\text{U}(1)$ is Abelian to get

$$\mathcal{A}_k = e^{-i\chi(\boldsymbol{\lambda})} \mathcal{A}_j e^{i\chi(\boldsymbol{\lambda})} + e^{-i\chi(\boldsymbol{\lambda})} d e^{i\chi(\boldsymbol{\lambda})} = \mathcal{A}_j + i \partial_{\boldsymbol{\lambda}} \chi(\boldsymbol{\lambda}) d\boldsymbol{\lambda}, \quad (4.37)$$

replicating Eq. (4.20) upon identifying $\mathcal{A} = \langle n(\boldsymbol{\lambda}) | d | n(\boldsymbol{\lambda}) \rangle = \langle n(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n(\boldsymbol{\lambda}) \rangle d\boldsymbol{\lambda}$. In terms of the Berry connection $\mathcal{A}$, the geometric phase in Eq. (4.14) takes the form

$$\eta(t) = i \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \langle n(\boldsymbol{\lambda}) | d | n(\boldsymbol{\lambda}) \rangle = i \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \mathcal{A}_{\mu}(\boldsymbol{\lambda}) d\lambda_{\mu}, \quad (4.38)$$

where we wrote

$$\langle n(\boldsymbol{\lambda}) | d | n(\boldsymbol{\lambda}) \rangle = \langle n(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n(\boldsymbol{\lambda}) \rangle d\boldsymbol{\lambda} = \langle n(\boldsymbol{\lambda}) | \partial_{\lambda_{\mu}} | n(\boldsymbol{\lambda}) \rangle d\lambda_{\mu} = \mathcal{A}_{\mu}(\boldsymbol{\lambda}) d\lambda_{\mu}. \quad (4.39)$$

With this, the geometric phase factor from Eq. (4.24) becomes

$$u(t) = \exp[i\eta(t)] = \exp \left[- \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \mathcal{A}_{\mu}(\boldsymbol{\lambda}) d\lambda_{\mu} \right]. \quad (4.40)$$

A comparison between Eq. (4.40) and the general expression for parallel transport in Eq. (2.166) allows us to identify the geometric phase factor as parallel transport in the $\text{U}(1)$ state bundle, governed by the Berry connection. The Berry phase then corresponds to parallel transport along closed curves, $C \subset \Lambda$, and thus to an element of the holonomy group $\text{Hol}(\mathcal{A})$ of the Berry connection $\mathcal{A}$. Based on the previous discussion of holonomy, we can predict that non-trivial Berry phases may arise whenever the parameter manifold Λ is not simply connected, i.e. when $\pi_1(\Lambda) \neq \{e\}$, or the Berry connection is non-flat, i.e. when the Berry curvature two-form $\mathcal{F} \equiv d\mathcal{A} = \mathcal{F}_{\mu\nu}(\boldsymbol{\lambda}) d\lambda_{\mu} \wedge d\lambda_{\nu}$ is not vanishing.

It is worth noting that the physics and mathematics literature frequently adopt different conventions when defining the Berry connection. In the physics literature, the Berry connection is often written as

$$\mathcal{A}^{\text{P}} \equiv i \langle n(\boldsymbol{\lambda}) | d | n(\boldsymbol{\lambda}) \rangle \in \mathbb{R}, \quad (4.41)$$

so that the Berry phase factor from Eq. (4.24) takes the form

$$u(C) = \exp \left[- \int_C \langle n(\boldsymbol{\lambda}) | d | n(\boldsymbol{\lambda}) \rangle \right] = \exp \left[i \int_C \mathcal{A}^{\text{P}} \right]. \quad (4.42)$$

With $u(C) \equiv \exp[i\eta(C)]$ from Eq. (4.24), the Berry phase $\eta(C)$ then becomes

$$\eta(C) = \int_C \mathcal{A}^{\text{P}}. \quad (4.43)$$

In contrast, the mathematical literature prefers the definition

$$\mathcal{A}^{\text{m}} \equiv \langle n(\boldsymbol{\lambda}) | d | n(\boldsymbol{\lambda}) \rangle \in i\mathbb{R}, \quad (4.44)$$

with which Eq. (4.24) takes the form

$$u(C) = \exp \left[- \int_C \langle n(\boldsymbol{\lambda}) | d | n(\boldsymbol{\lambda}) \rangle \right] = \exp \left[- \int_C \mathcal{A}^{\text{m}} \right], \quad (4.45)$$

and the Berry phase becomes

$$\eta(C) = i \int_C \mathcal{A}^{\text{m}}, \quad (4.46)$$

as in Eq. (4.14). Comparing Eqs. (4.45) and (4.40) shows that we have adopted the mathematical definition of the Berry connection above. While we will typically follow this convention to ease the mathematical interpretation, both definitions are equivalent and may be used interchangeably.

4.3 The Wilczek–Zee Phase

The Berry phase is limited to adiabatic evolutions of isolated states, i.e. states that are non-degenerate everywhere on the parameter manifold. A natural generalisation of Berry’s adiabatic phase arises when we consider degenerate subspaces of instantaneous eigenstates. For this part, we follow Ref. [80].

Let $H(\boldsymbol{\lambda})$ be a parameter-dependent Hamiltonian with a degenerate spectrum of d_n -fold degenerate eigenvalues $E_n(\boldsymbol{\lambda})$. Suppose that the degeneracy d_n of the eigenspaces

$$\mathcal{H}_n(\boldsymbol{\lambda}) = \text{span}(|n_1(\boldsymbol{\lambda})\rangle, \dots, |n_{d_n}(\boldsymbol{\lambda})\rangle) \quad (4.47)$$

is parameter-independent, i.e. that $\dim_{\mathbb{C}}(\mathcal{H}_n) = d_n$ for all parameter configurations $\boldsymbol{\lambda}$ of the parameter manifold Λ . In particular, we assume that there is no level-crossing between the eigenvalues $E_n(\boldsymbol{\lambda})$ and $E_m(\boldsymbol{\lambda})$, so that degenerate subspaces $\mathcal{H}_n(\boldsymbol{\lambda})$ and $\mathcal{H}_m(\boldsymbol{\lambda})$ do not intersect for $n \neq m$. Let us select a degenerate eigenvalue $E_n(\boldsymbol{\lambda})$ with $d_n > 1$ and prepare the system in a generic initial state

$$|\psi(0)\rangle = \sum_{j=1}^{d_n} c_j(0) |n_j(\boldsymbol{\lambda}(0))\rangle. \quad (4.48)$$

If the parameters $\boldsymbol{\lambda} = \boldsymbol{\lambda}(t)$ change slowly as a function of time, the adiabatic theorem still ensures that the system remains in the eigenspace $\mathcal{H}_n$ but now there is more than one state to account for within this eigenspace. To do this, we choose a generic ansatz

$$|\psi(t)\rangle = \sum_{j=1}^{d_n} c_j(t) |n_j(\boldsymbol{\lambda}(t))\rangle. \quad (4.49)$$

With this, the time-dependent Schrödinger equation from Eq. (4.3) reads

$$i \sum_{j=1}^{d_n} \left(\frac{dc_j(t)}{dt} |n_j(\boldsymbol{\lambda}(t))\rangle + c_j(t) \frac{d}{dt} |n_j(\boldsymbol{\lambda}(t))\rangle \right) = E_n(\boldsymbol{\lambda}(t)) \sum_{j=1}^{d_n} c_j(t) |n_j(\boldsymbol{\lambda}(t))\rangle, \quad (4.50)$$

which reduces to a set of differential equations

$$\frac{dc_j(t)}{dt} = - \sum_{k=1}^{d_n} \left(i E_n(\boldsymbol{\lambda}(t)) \delta_{jk} + \langle n_j(\boldsymbol{\lambda}(t)) | \frac{d}{dt} | n_k(\boldsymbol{\lambda}(t)) \rangle \right) c_k(t) \quad (4.51)$$

for the coefficients $c_j(t)$. The solutions of Eq. (4.51) are given by

$$\begin{aligned} c_j(t) &= \sum_{k=1}^{d_n} \left(\mathcal{T} \exp \left[- \int_0^t \left(i E_n(\boldsymbol{\lambda}(\tau)) \mathbb{1}_{d_n} + \mathcal{A}_{d_n}(\boldsymbol{\lambda}(\tau)) \right) d\tau \right] \right)_{jk} c_k(0) \\ &\stackrel{(\diamond)}{=} \exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau \right] \sum_{k=1}^{d_n} \left(\mathcal{P} \exp \left[- \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \mathcal{A}_{d_n}(\boldsymbol{\lambda}(\tau)) d\boldsymbol{\lambda} \right] \right)_{jk} c_k(0), \end{aligned} \quad (4.52)$$

where $\mathcal{T}$ and $\mathcal{P}$ denote the time-ordering and the path-ordering operators, and where $\mathbb{1}_{d_n}$ represents the $d_n \times d_n$ identity matrix. In $(\diamond)$, we used that $\exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) \mathbb{1}_{d_n} d\tau \right]$ commutes with everything and merely acts as a global phase factor. Moreover, we have identified the non-Abelian Berry connection $\mathcal{A}_{d_n}$ of the n -th degenerate eigenspace $\mathcal{H}_n$. It is a skew-Hermitian $d_n \times d_n$ matrix defined in terms of its elements

$$\mathcal{A}_{d_n, jk}(\boldsymbol{\lambda}) = \langle n_j(\boldsymbol{\lambda}(\tau)) | \frac{d}{d\tau} | n_k(\boldsymbol{\lambda}(\tau)) \rangle d\tau = \langle n_j(\boldsymbol{\lambda}) | \partial_{\lambda_\mu} | n_k(\boldsymbol{\lambda}) \rangle d\lambda_\mu. \quad (4.53)$$

If we substitute Eq. (4.52) into Eq. (4.49) we obtain

$$|\psi(t)\rangle = \sum_{j,k}^{d_n} |n_j(\boldsymbol{\lambda}(t))\rangle \exp \left[-i \int_0^t E_n(\boldsymbol{\lambda}(\tau)) d\tau \right] \left(\mathcal{P} \exp \left[- \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \mathcal{A}_{d_n}(\boldsymbol{\lambda}) d\boldsymbol{\lambda} \right] \right)_{jk} c_k(0). \quad (4.54)$$

Note that $\mathcal{A}_{d_n}(\boldsymbol{\lambda})$ is called *non-Abelian* because it is a one-form with values in the skew-Hermitian $d_n \times d_n$ matrices. For $d_n > 1$ these matrices do not commute and the path ordered matrix exponential

$$\mathcal{U}_{d_n}(P) = \mathcal{P} \exp \left[- \int_{\boldsymbol{\lambda}(0)}^{\boldsymbol{\lambda}(t)} \mathcal{A}_{d_n}(\boldsymbol{\lambda}) \right] = \mathcal{P} \exp \left[- \int_P \mathcal{A}_{d_n}(\boldsymbol{\lambda}) \right] \quad (4.55)$$

defines an element of the non-Abelian unitary group $U(d_n)$ that depends only on the geometric properties of the path $P : [0, t] \rightarrow \Lambda$, $\tau \mapsto \boldsymbol{\lambda}(\tau)$ in the parameter manifold. The $U(d_n)$ matrix

$$\mathcal{U}_{d_n}(C) = \mathcal{P} \exp \left[- \oint_C \mathcal{A}_{d_n}(\boldsymbol{\lambda}) \right] \equiv \mathcal{P} \exp [i\mathcal{W}_{d_n}(C)] \quad (4.56)$$

that arises after traversing a closed curve $C : [0, t] \rightarrow \Lambda$, $\tau \mapsto \boldsymbol{\lambda}(\tau)$ with $\boldsymbol{\lambda}(0) = \boldsymbol{\lambda}(t)$ is usually called the Wilczek–Zee phase factor and constitutes a generalisation of the $U(1)$ Berry phase [39, 80]. The Hermitian $d_n \times d_n$ matrix

$$\mathcal{W}_{d_n}(C) = i \oint_C \mathcal{A}_{d_n}(\boldsymbol{\lambda}) \quad (4.57)$$

is called the Wilczek–Zee phase and readily generalises Eq. (4.14). Unlike the Berry phase factor, which is invariant under gauge transformations, the Wilczek–Zee phase factor is only gauge covariant. This is an immediate consequence of $\dim_{\mathbb{C}}(\mathcal{H}_n) > 1$, due to which the generic gauge transformations among the eigenstates spanning $\mathcal{H}_n(\boldsymbol{\lambda})$ take the form

$$|n'_j(\boldsymbol{\lambda})\rangle = \sum_k |n_k(\boldsymbol{\lambda})\rangle U_{kj}(\boldsymbol{\lambda}), \quad (4.58)$$

where $U(\boldsymbol{\lambda}) \in U(d_n)$ implements some unitary rotation of eigenstates within $\mathcal{H}_n(\boldsymbol{\lambda})$. Repeating the calculation in Eq. (4.20) yields the transformation behaviour

$$\begin{aligned} \mathcal{A}'_{d_n}(\boldsymbol{\lambda})_{jk} &= \langle n'_j(\boldsymbol{\lambda}) | \partial_{\lambda_\mu} | n'_k(\boldsymbol{\lambda}) \rangle d\lambda_\mu \\ &= \sum_{r,s=1}^{d_n} U_{jr}^\dagger(\boldsymbol{\lambda}) \langle n_r(\boldsymbol{\lambda}) | \partial_{\lambda_\mu} | n_s(\boldsymbol{\lambda}) \rangle U_{sk}(\boldsymbol{\lambda}) d\lambda_\mu \\ &= \sum_{r,s=1}^{d_n} U_{jr}^\dagger(\boldsymbol{\lambda}) \langle n_r(\boldsymbol{\lambda}) | \left[(\partial_{\lambda_\mu} | n_s(\boldsymbol{\lambda}) \rangle) U_{sk}(\boldsymbol{\lambda}) + | n_s(\boldsymbol{\lambda}) \rangle (\partial_{\lambda_\mu} U_{sk}(\boldsymbol{\lambda})) \right] d\lambda_\mu \\ &= \sum_{r,s=1}^{d_n} \left[U_{jr}^\dagger(\boldsymbol{\lambda}) \langle n_r(\boldsymbol{\lambda}) | (\partial_{\lambda_\mu} | n_s(\boldsymbol{\lambda}) \rangle) U_{sk}(\boldsymbol{\lambda}) + U_{jr}^\dagger(\boldsymbol{\lambda}) \langle n_r(\boldsymbol{\lambda}) | n_s(\boldsymbol{\lambda}) \rangle (\partial_{\lambda_\mu} U_{sk}(\boldsymbol{\lambda})) \right] d\lambda_\mu \\ &= \sum_{r,s=1}^{d_n} \left[U_{jr}^\dagger(\boldsymbol{\lambda}) \langle n_r(\boldsymbol{\lambda}) | (\partial_{\lambda_\mu} | n_s(\boldsymbol{\lambda}) \rangle) U_{sk}(\boldsymbol{\lambda}) + U_{jr}^\dagger(\boldsymbol{\lambda}) \delta_{rs} (\partial_{\lambda_\mu} U_{sk}(\boldsymbol{\lambda})) \right] d\lambda_\mu \\ &= \sum_{r,s=1}^{d_n} U_{jr}^\dagger(\boldsymbol{\lambda}) \mathcal{A}_n(\boldsymbol{\lambda})_{rs} U_{sk}(\boldsymbol{\lambda}) + \sum_{r=1}^{d_n} U_{jr}^\dagger(\boldsymbol{\lambda}) (\partial_{\lambda_\mu} U(\boldsymbol{\lambda}))_{rk} d\lambda_\mu \\ &= (U^\dagger(\boldsymbol{\lambda}) \mathcal{A}_n(\boldsymbol{\lambda}) U(\boldsymbol{\lambda}))_{jk} + (U^\dagger(\boldsymbol{\lambda}) \partial_{\lambda_\mu} U(\boldsymbol{\lambda}))_{jk} d\lambda_\mu, \end{aligned} \quad (4.59)$$

which, in matrix notation, reads

$$\mathcal{A}'_{d_n}(\boldsymbol{\lambda}) = U^\dagger(\boldsymbol{\lambda}) \mathcal{A}_{d_n}(\boldsymbol{\lambda}) U(\boldsymbol{\lambda}) + U^\dagger(\boldsymbol{\lambda}) \partial_{\lambda_\mu} U(\boldsymbol{\lambda}) d\lambda_\mu. \quad (4.60)$$

Since $\mathcal{A}_{d_n}(\boldsymbol{\lambda})$ and $U(\boldsymbol{\lambda})$ do not generally commute, this expression cannot be further simplified. If we plug Eq. (4.60) into the definition of the Wilczek–Zee phase factor given in Eq. (4.56), we find that it transforms as

$$\mathcal{U}'_{d_n}(C) = U^\dagger(C) \mathcal{U}_{d_n}(C) U(C), \quad (4.61)$$

so that $\mathcal{U}_{d_n}(C)$ turns out gauge covariant rather than gauge invariant. As a result, the Wilczek–Zee phase factor is typically not observable and one has to rely on related gauge invariant quantities like the Wilson loop

$$w(C) = \text{tr}(\mathcal{U}_{d_n}(C)) \quad (4.62)$$

of $\mathcal{U}_{d_n}(C)$, which is manifestly gauge invariant due to the cyclic property of the trace.

Just like in the Abelian case, we can translate the above considerations to the language of fibre bundles. Specifically, the parameter-dependent eigenspaces $\mathcal{H}_n(\boldsymbol{\lambda})$ define a principal $U(d_n)$ -bundle $P \xrightarrow{\pi} \Lambda$ over the parameter manifold Λ . The non-Abelian Berry connection $\mathcal{A}_n$ with elements

$$\mathcal{A}_{d_n,jk} = \langle n_j(\boldsymbol{\lambda}) | d | n_k(\boldsymbol{\lambda}) \rangle \quad (4.63)$$

defines an Ehresmann connection. As before, this follows from its skew-Hermiticity

$$\begin{aligned} 0 &= d(\langle n_j(\boldsymbol{\lambda}) | n_k(\boldsymbol{\lambda}) \rangle) = (d \langle n_j(\boldsymbol{\lambda}) |) | n_k(\boldsymbol{\lambda}) \rangle + \langle n_j(\boldsymbol{\lambda}) | d | n_k(\boldsymbol{\lambda}) \rangle \\ &= \langle n_k(\boldsymbol{\lambda}) | d | n_j(\boldsymbol{\lambda}) \rangle^* + \langle n_j(\boldsymbol{\lambda}) | d | n_k(\boldsymbol{\lambda}) \rangle \\ &= \mathcal{A}_{d_n,kj}^* + \mathcal{A}_{d_n,jk} \\ &= \mathcal{A}_{d_n,jk}^\dagger + \mathcal{A}_{d_n,jk} , \end{aligned} \quad (4.64)$$

which tells us that $\mathcal{A}_{d_n}$ takes values in the Lie algebra $\mathfrak{u}(d_n)$ of $U(d_n)$, and the transformation behaviour from Eq. (4.60), which replicates the transformation behaviour Eq. (2.154) of an Ehresmann connection. The unitary matrix $\mathcal{U}_{d_n}(T)$ from Eq. (4.55) then corresponds to the parallel transport specified by the non-Abelian Berry connection $\mathcal{A}_{d_n}$, and the Wilczek–Zee phase factor $\mathcal{U}_{d_n}(C)$ from Eq. (4.56) captures its $U(d_n)$ holonomy. For the sake of completeness, we note that the non-Abelian Berry curvature $\mathcal{F}_{d_n}$ takes the form

$$\mathcal{F}_{d_n} = d\mathcal{A}_{d_n} + \mathcal{A}_{d_n} \wedge \mathcal{A}_{d_n} , \quad (4.65)$$

as specified in Eq. (2.188). In contrast to the Abelian case, the wedge product $\mathcal{A}_{d_n} \wedge \mathcal{A}_{d_n}$ is generally not equal to zero because the commutator $[\mathcal{A}_{d_n,\mu}, \mathcal{A}_{d_n,\nu}]$ between the $\mathfrak{u}(d_n)$ -valued components $\mathcal{A}_{d_n,\mu}$ of $\mathcal{A}_{d_n} = \mathcal{A}_{d_n,\mu} d\lambda_\mu$ is usually non-trivial.

4.4 Berry, Aharonov and Anandan

Berry’s interpretation of the geometric phase is based on an adiabatic constraint for parameter-dependent families of Hamiltonians $H(\boldsymbol{\lambda})$, which ensures that the time-evolved quantum state $|\psi(t)\rangle$ of the system $H(\boldsymbol{\lambda})$ clings to the d_n -dimensional instantaneous eigenspace $\mathcal{H}_n(\boldsymbol{\lambda})$ if $|\psi(0)\rangle \in \mathcal{H}_n(\boldsymbol{\lambda}(0))$. As a result, the adiabatic time evolution $|\psi(t)\rangle$ acquires a purely geometric contribution that can be understood in terms of parallel transport in the vector bundle of instantaneous eigenspaces $\mathcal{H}_n(\boldsymbol{\lambda})$ or, equivalently, the associated principal $U(d_n)$ -bundle over the parameter manifold Λ . Both the Berry phase and the Wilczek–Zee phase arise from this description for $d_n = 1$ and $d_n > 0$ respectively.

In Ref. [79], Aharonov and Anandan present an alternative perspective that allows for a generalisation of this concept beyond the adiabatic regime. They propose to shift the focus away from the cyclic evolution of Hamiltonians and towards the direct cyclic evolution of quantum states. The central idea is that every cyclic evolution of quantum states $|\psi(t)\rangle \in \mathcal{H}$ defines a closed curve in the projective Hilbert space $P(\mathcal{H})$ of $\mathcal{H}$. To see this, recall that physical states are equivalence classes $[|\psi\rangle]$ of rays, as given in Eq. (4.1). An evolution $|\psi(0)\rangle \rightarrow |\psi(T)\rangle$ of state representatives is therefore cyclic if

$$|\psi(T)\rangle = c |\psi(0)\rangle \quad (4.66)$$

for some $c \in \mathbb{C}$, since then

$$[|\psi(T)\rangle] = [|\psi(0)\rangle] . \quad (4.67)$$

Moreover, the projective Hilbert space $P(\mathcal{H})$ is precisely the complex manifold formed by the equivalence classes of quantum states. For a d -dimensional Hilbert space $\mathcal{H}$ with $d < \infty$, it can be constructed as the quotient

$$P(\mathcal{H}) = \mathbb{S}^{2d-1}/U(1) \simeq \mathbb{CP}^{d-1} \quad (4.68)$$

of normalised states $\{|\psi\rangle \in \mathcal{H}, \|\psi\|^2 = 1\} \simeq \mathbb{S}^{2d-1}$ modulo phase factors $e^{i\varphi} \in U(1)$. As a consequence, every cyclic evolution $|\psi(0)\rangle \rightarrow |\psi(T)\rangle = c|\psi(0)\rangle$ defines a closed curve

$$\mathcal{C} : [0, T] \rightarrow P(\mathcal{H}), \quad t \mapsto [|\psi(t)\rangle] \equiv |\psi(t)\rangle \langle \psi(t)| \quad (4.69)$$

in the projective Hilbert space $P(\mathcal{H})$. For later reference, we have included the rather common notation $[|\psi(t)\rangle] \equiv |\psi(t)\rangle \langle \psi(t)|$, where equivalence classes $[|\psi(t)\rangle]$ of physical states are represented by manifestly gauge-independent projectors $|\psi(t)\rangle \langle \psi(t)|$. Note that apart from the cyclicity condition there are no constraints for $|\psi(t)\rangle$. In particular, it need not be an energy eigenstate.

So what are the implications of this? Following Ref. [80], we consider a time-dependent Hamiltonian $H(\boldsymbol{\lambda}(t))$, whose time-dependence is determined by some (not necessarily closed) curve

$$\Gamma : [0, T] \rightarrow \Lambda, \quad t \mapsto \boldsymbol{\lambda}(t) \quad (4.70)$$

in a parameter manifold Λ . The time evolution $|\psi(t)\rangle$ of every quantum state $|\psi\rangle \in \mathcal{H}$ defines a curve

$$\Psi : [0, T] \rightarrow P(\mathcal{H}), \quad t \mapsto [|\psi(t)\rangle] \quad (4.71)$$

in the projective Hilbert space. In particular, every *cyclic* state defines a *closed* curve. The time evolution of any cyclic state $|\psi\rangle$ must be of the form

$$|\psi(T)\rangle = \exp[-i\alpha(T)] |\psi(0)\rangle. \quad (4.72)$$

Using an argument analogous to that used in the derivation of the Berry phase, one can show that the overall phase $\alpha(T)$ separates as

$$\alpha(T) = \gamma(T) - \zeta(\mathcal{C}) \quad (4.73)$$

into a dynamical phase $\gamma(T)$ and a geometric phase $\zeta(\mathcal{C})$, which depends only on the closed curve $\mathcal{C} \subset P(\mathcal{H})$ traced by $|\psi(t)\rangle$. Since $|\psi(t)\rangle$ is not necessarily an energy eigenstate, the dynamical phase takes the form

$$\gamma(T) = \int_0^T \langle \psi(\tau) | H(\boldsymbol{\lambda}(\tau)) | \psi(\tau) \rangle d\tau. \quad (4.74)$$

Meanwhile, the geometric phase is given by

$$\zeta(\mathcal{C}) = i \int_0^T \langle \phi(\tau) | \partial_\tau | \phi(\tau) \rangle dt = i \oint_{\mathcal{C}} \langle \phi(\tau) | d_P | \phi(\tau) \rangle, \quad (4.75)$$

where $|\phi(\tau)\rangle$ is defined via a section

$$\sigma : P(\mathcal{H}) \rightarrow E, \quad [|\psi\rangle] \mapsto \sigma([|\psi\rangle]) := |\phi\rangle \quad (4.76)$$

of the so-called *tautological* line bundle $E \xrightarrow{\pi} P(\mathcal{H})$, whose fibre $F_{[|\psi\rangle]}$ over every point $[|\psi\rangle] \in P(\mathcal{H})$ is exactly the one-dimensional subspace of $\mathcal{H}$ spanned by $|\psi\rangle$, i.e.

$$F_{[|\psi\rangle]} = \{c|\psi\rangle, c \in \mathbb{C}^*\}. \quad (4.77)$$

Note that the total differential d_P in Eq. (4.75) is the total differential of $P(\mathcal{H})$ rather than that of the parameter manifold Λ . One can define the Aharonov–Anandan connection

$$\mathcal{A}_{AA} := \langle \phi(\tau) | d_P | \phi(\tau) \rangle \quad (4.78)$$

to show that the Aharonov–Anandan phase in Eq. (4.75) corresponds to the $U(1)$ holonomy of the tautological line bundle $E \xrightarrow{\pi} P(\mathcal{H})$. This holonomy is often discussed in terms of the so-called Aharonov–Anandan principal $U(1)$ -bundle $P_{AA} \xrightarrow{\pi_{AA}} P(\mathcal{H})$, which is precisely the principal $U(1)$ -bundle associated to $E \xrightarrow{\pi} P(\mathcal{H})$. Note that the Aharonov–Anandan phase replicates the Berry phase if the parameter evolution $\boldsymbol{\lambda}(t)$ is closed, i.e. $\boldsymbol{\lambda}(0) = \boldsymbol{\lambda}(T)$, and $[\psi(t)] = [n(\boldsymbol{\lambda}(t))]$ parameterises a curve in $P(\mathcal{H})$ that corresponds to the n -th eigenstate of $H(\boldsymbol{\lambda}(t))$. In this sense, the Aharonov–Anandan phase constitutes a generalisation of the Berry phase.

Just like the Berry phase, the Aharonov–Anandan phase can be generalised to δ -dimensional cyclic subspaces $\mathcal{S}(t) \subset \mathcal{H}$. The difference is that these subspaces need not be eigenspaces of $H(\boldsymbol{\lambda}(t))$; they only have to be cyclic under a given parameter evolution $\boldsymbol{\lambda}(t)$. Analogous to how a cyclic state $|\psi(t)\rangle$ defines a closed curve in the projective Hilbert space $P(\mathcal{H})$ of physical states, a cyclic δ -dimensional subspace $\mathcal{S}(t) \simeq \mathbb{C}^\delta$ defines a closed curve

$$\mathcal{C} : [0, T] \rightarrow G_\delta(\mathcal{H}), \quad t \mapsto [\mathcal{S}(t)], \quad (4.79)$$

in the complex manifold $G_\delta(\mathcal{H})$ of equivalence classes $[\mathcal{S}]$ of δ -dimensional subspaces $\mathcal{S} \subset \mathcal{H}$, which is known as the Grassmannian.³ Here, a cyclic subspace $\mathcal{S}(t)$ is any subspace that returns to itself, fulfilling $\mathcal{S}(0) \rightarrow \mathcal{S}(T) = \mathcal{S}(0)$. In the projector notation, this means that $P_{\mathcal{S}(0)} \rightarrow P_{\mathcal{S}(T)} = P_{\mathcal{S}(0)}$ throughout the cyclic evolution. If we choose to spell out

$$P_{\mathcal{S}(0)} = \sum_{j=1}^{\delta} |\phi_j(\mathcal{S}(0))\rangle \langle \phi_j(\mathcal{S}(0))| \quad (4.80)$$

in terms of an arbitrary but fixed δ -frame

$$f(P_{\mathcal{S}(0)}) = \left\{ |\phi_1(\mathcal{S}(0))\rangle, \dots, |\phi_\delta(\mathcal{S}(0))\rangle \right\} \quad (4.81)$$

of δ orthonormal basis vectors $|\phi_j(\mathcal{S}(0))\rangle$ of $\mathcal{S}(0)$, then the cyclicity of $P_{\mathcal{S}(0)}$ tells us that the $|\phi_j(\mathcal{S}(T))\rangle$ can at most transform as

$$|\phi_j(\mathcal{S}(T))\rangle = \sum_{k=1}^{\delta} |\phi_k(\mathcal{S}(0))\rangle \mathcal{U}_{kj}(\mathcal{S}(T)), \quad (4.82)$$

with a unitary $\delta \times \delta$ transformation matrix $\mathcal{U}(\mathcal{S}(T)) \in U(\delta)$. Thus, it is not surprising that the general time evolution of an initial state

$$|\psi(0)\rangle = \sum_{k=1}^{\delta} |\phi_k(\mathcal{S}(0))\rangle c_k(0) \in \mathcal{S}(0) \quad (4.83)$$

takes the form

$$|\psi(t)\rangle = \sum_{j,k}^{\delta} |\phi_j(\mathcal{S}(0))\rangle \mathcal{U}_{jk}(\mathcal{S}(t)) c_k(0), \quad (4.84)$$

where $\mathcal{U}_{jk}(\mathcal{S}(t))$ are the elements of a unitary $\delta \times \delta$ matrix

$$\mathcal{U}(\mathcal{S}(t)) := \mathcal{T} \exp \left[- \int_0^t \left(i\mathcal{E}(\mathcal{S}(\tau)) d\tau + \mathcal{A}_\delta(\mathcal{S}(\tau)) \right) \right], \quad (4.85)$$

which is defined in terms of $\delta \times \delta$ Hermitian matrices $\mathcal{E}(\mathcal{S}(\tau))$ and $\mathcal{A}_\delta(\mathcal{S}(\tau))$ with elements

$$\mathcal{E}(\mathcal{S}(\tau))_{jk} := \langle \phi_j(\mathcal{S}(\tau)) | H(\boldsymbol{\lambda}(\tau)) | \phi_k(\mathcal{S}(\tau)) \rangle \quad \text{and} \quad \mathcal{A}_\delta(\mathcal{S}(\tau))_{jk} := \langle \phi_j(\mathcal{S}(\tau)) | \frac{d}{d\tau} | \phi_k(\mathcal{S}(\tau)) \rangle d\tau, \quad (4.86)$$

³We have encountered the Grassmannian $G_n(\mathbb{C}^\infty)$ as the classifying space of n -dimensional complex vector bundles and their associated $U(n)$ -principal bundles in Sec. 2.2.3. The Grassmannian is the natural generalisation of the projective Hilbert space to equivalence classes of higher-dimensional subspaces.

where $\mathcal{E}(\mathcal{S}(t))$ is the generalised non-Abelian dynamical phase, and $\mathcal{A}_\delta(\mathcal{S}(t))$ is the non-Abelian Aharonov–Anandan connection of the tautological frame bundle $E \xrightarrow{\pi_{AA}} G_\delta(\mathcal{H})$ of δ -frames over the Grassmannian $G_\delta(\mathcal{H})$. Again, Eq. (4.84) is the non-adiabatic generalisation of Eq. (4.54) and reproduces it in the adiabatic limit with degenerate eigenspaces $\mathcal{S}(t) = \mathcal{H}_n(\lambda(t))$. The important difference is that the matrices $\mathcal{E}(\mathcal{S}(\tau))$ and $\mathcal{A}_\delta(\mathcal{S}(\tau))$ from Eq. (4.85) do not generally commute. As a consequence, it is generally not possible to write $\mathcal{U}(\mathcal{S}(t))$ as a product

$$\mathcal{U}(\mathcal{S}(t)) = \mathcal{U}_{\text{dyn.}}(\mathcal{S}(t)) \mathcal{U}_{\text{geom.}}(\mathcal{S}(t)) \quad (4.87)$$

of the generalised dynamical phase matrix

$$\mathcal{U}_{\text{dyn.}}(\mathcal{S}(t)) := \mathcal{T} \exp \left[-i \int_0^t \mathcal{E}(\mathcal{S}(\tau)) d\tau \right], \quad (4.88)$$

and the generalised geometric phase matrix

$$\mathcal{U}_{\text{geom.}}(\mathcal{C}) := \mathcal{T} \exp \left[- \int_0^t \mathcal{A}_\delta(\mathcal{S}(\tau)) \right] = \mathcal{P} \exp \left[- \oint_{\mathcal{C}} \mathcal{A}_\delta(\mathcal{S}) \right]. \quad (4.89)$$

Even though their product does not generally determine the unitary time evolution in Eq. (4.84), the two unitary matrices $\mathcal{U}_{\text{dyn.}}(\mathcal{S}(t))$ and $\mathcal{U}_{\text{geom.}}(\mathcal{S}(t))$ are still defined independently. In fact, the non-Abelian Aharonov–Anandan phase factor $\mathcal{U}_{\text{geom.}}(\mathcal{S}(t))$ still captures the $U(\delta)$ holonomy of the tautological frame bundle $E \xrightarrow{\pi_{AA}} G_\delta(\mathcal{H})$. As before, this holonomy is often discussed in terms of the non-Abelian Aharonov–Anandan principal $U(\delta)$ -bundle $P_{AA} \xrightarrow{\pi_{AA}} G_\delta(\mathcal{H})$, which is precisely the associated principal $U(\delta)$ -bundle to the tautological frame bundle $E \xrightarrow{\pi} G_\delta(\mathcal{H})$. Just like Eq. (4.56), the geometric phase factor in Eq. (4.89) is only gauge covariant under gauge transformations of the δ -frames $f(P_{\mathcal{S}(t)})$ so the only observables associated to $\mathcal{U}_{\text{geom.}}(\mathcal{S}(t))$ are quantities like its Wilson loop

$$w[\mathcal{C}] = \text{tr}(\mathcal{U}_{\text{geom.}}(\mathcal{C})). \quad (4.90)$$

4.5 Geometric Features of Diabatic Dynamics

In previous sections, we have discussed geometric phases in the context of adiabatic dynamics: Berry’s Abelian phase for non-degenerate eigenstates, and the non-Abelian Wilczek–Zee phase for degenerate eigenspaces. We also addressed the Aharonov–Anandan generalisation, which extends these concepts beyond the adiabatic limit and captures the geometry of cyclic subspaces that need not be eigenspaces of a given parameter-dependent Hamiltonian.

There is a third regime that is of practical importance for physics: *diabatic dynamics* of systems with *slightly non-degenerate spectra*, where small energy splittings arise, for example, due to proximity-induced hybridisation between eigenstates with once-degenerate energies. Consider a system with a d_n -dimensional almost-degenerate eigenspace $\mathcal{H}_n(\lambda(t))$ spanned by energy eigenstates, whose energies lie within a narrow range $\Delta\epsilon$, but remain separated by at least $\Delta E \gg \Delta\epsilon$ from the rest of the spectrum. Suppose this system evolves on a time scale θ satisfying

$$\frac{\hbar}{\Delta E} \ll \theta \lesssim \frac{\hbar}{\Delta\epsilon}. \quad (4.91)$$

While the separation $\hbar/\Delta E \ll \theta$ of time scales ensures that transitions out of $\mathcal{S}(t)$ are negligible and that states $|\psi\rangle \in \mathcal{H}_n(\lambda(t))$ evolve approximately cyclically under periodic parameter variations, the similarity $\theta \lesssim \hbar/\Delta\epsilon$ implies that the dynamics within $\mathcal{H}_n(\lambda(t))$ are fast enough to induce mixing between the nearly-degenerate eigenstates. As a consequence, the evolution of states in $\mathcal{H}_n(\lambda(t))$ can no longer be described in terms of the original instantaneous eigenstates, and requires the non-Abelian non-adiabatic Aharonov–Anandan description from Eq. (4.84) instead. In particular, the dynamical and geometric contributions $\mathcal{E}$ and $\mathcal{A}$ to Eq. (4.85) do not commute and are inseparably intertwined in the resulting unitary evolution matrix $\mathcal{U}$.

In such situations, *diabatic* dynamics can be used to probe the geometric evolution of the subspace $\mathcal{H}_n(\boldsymbol{\lambda}(t))$ even in the presence of finite energy splittings. Specifically, if the system dynamics can be accelerated such that typical time scale θ satisfies

$$\frac{\hbar}{\Delta E} \ll \theta \ll \frac{\hbar}{\Delta \epsilon}, \quad (4.92)$$

the dynamics become too fast to resolve the energy splittings within $\mathcal{H}_0(\phi)$ and their degeneracy is effectively restored. If this is possible, the non-Abelian part of the unitary evolution $\mathcal{U}$ is once more effectively determined by the holonomy of the principal $U(d_n)$ -bundle $P \xrightarrow{\pi} \Lambda$ associated with the instantaneous eigenspaces $\mathcal{H}_n(\boldsymbol{\lambda}(t))$, and the formalism by Wilczek and Zee remains applicable to leading order, despite the lack of exact degeneracy.

Note that Eq. (4.92) requires a careful balance between the adiabatic and diabatic time scale separations. This can be very difficult in practice. In particular, it may not always be possible to assume time scales that are short enough to render the system insensitive to the small splittings, but long enough to suppress excitations beyond the nearly degenerate subspace.

4.6 Discretised Geometric Phases

While the geometric phases introduced above are generally defined for continuous (or even smooth) bundles, their practical computation often relies on numerical methods that require discretised versions. In the following, we briefly outline discretisations of the Abelian (Berry) and non-Abelian (Wilczek–Zee) geometric phases for later use. Based on these, we are going to evaluate geometric phases for discrete families of quantum states numerically. Concretely, we will apply them to non-degenerate BdG ground states to determine Sombrero Berry phases in Chap. 9, and to nearly-degenerate low-energy BdG states to study exchangeless braiding phenomena in Chap. 10.

The discretised expressions for the Berry and WZ phases will be derived from Eq. (4.56), with the (Abelian) Berry phase appearing as the special case $d_n = 1$. As a starting point, we discretise the closed curve

$$C : [0, t] \rightarrow \Lambda, \quad \tau \mapsto C(\tau) \equiv \boldsymbol{\lambda}(\tau), \quad (4.93)$$

with $\boldsymbol{\lambda}(0) = \boldsymbol{\lambda}(t)$ as

$$C_I \equiv \{\boldsymbol{\lambda}_0, \boldsymbol{\lambda}_1, \dots, \boldsymbol{\lambda}_{I-1}, \boldsymbol{\lambda}_I = \boldsymbol{\lambda}_0\} \subset C, \quad (4.94)$$

where the parameter $I \in \mathbb{N}$ controls the resolution of the discretisation. Specifically, we define

$$\boldsymbol{\lambda}_j = \boldsymbol{\lambda}(\tau_j), \quad (4.95)$$

using the discretised interval

$$[0, t]_I \equiv \{\tau_0 = 0, \tau_1, \dots, \tau_{I-1}, \tau_I = t\}. \quad (4.96)$$

While $[0, t]_I \subset [0, t]$ can in principle be any cardinality I subset, we typically choose a uniform partition

$$[0, t]_I \equiv \left\{0, \frac{t}{I}, \dots, \frac{(I-1)t}{I}, t\right\}, \quad (4.97)$$

where $\tau_j = (jt)/I$. Based on a uniform discretisation C_I , the Wilczek–Zee phase $\mathcal{W}_n(C)$ from Eq. (4.57) may be approximated by the (right) Riemann sum

$$\mathcal{W}_n(C) = i \oint_C \mathcal{A}_{d_n}(\boldsymbol{\lambda}) d\boldsymbol{\lambda} \approx i \sum_{\boldsymbol{\lambda}_j \in C_I} \mathcal{A}_{d_n}(\boldsymbol{\lambda}_j) \Delta \boldsymbol{\lambda}_j, \quad (4.98)$$

where $\Delta \boldsymbol{\lambda}_j \equiv (\boldsymbol{\lambda}_j - \boldsymbol{\lambda}_{j-1})$. Additionally, we can approximate the elements

$$\mathcal{A}_{d_n, kl}(\boldsymbol{\lambda}) = \langle n_k(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n_l(\boldsymbol{\lambda}) \rangle \quad (4.99)$$

of $\mathcal{A}_{d_n}(\boldsymbol{\lambda})$ using the (backward) difference quotient

$$\langle n_k(\boldsymbol{\lambda}) | \partial_{\boldsymbol{\lambda}} | n_l(\boldsymbol{\lambda}) \rangle \approx \frac{\langle n_k(\boldsymbol{\lambda}) | (| n_l(\boldsymbol{\lambda}) \rangle - | n_l(\boldsymbol{\lambda} - \Delta\boldsymbol{\lambda}) \rangle)}{\Delta\boldsymbol{\lambda}} = \frac{\delta_{kl} - \langle n_k(\boldsymbol{\lambda}) | n_l(\boldsymbol{\lambda} - \Delta\boldsymbol{\lambda}) \rangle}{\Delta\boldsymbol{\lambda}} \quad (4.100)$$

and introduce the overlap matrix $a_{d_n}(\boldsymbol{\lambda})$ with elements

$$a_{d_n,kl}(\boldsymbol{\lambda}) := \langle n_k(\boldsymbol{\lambda}) | n_l(\boldsymbol{\lambda} - \Delta\boldsymbol{\lambda}) \rangle \quad (4.101)$$

to write

$$\mathcal{A}_{d_n}(\boldsymbol{\lambda}) \approx \frac{\mathbb{1}_{d_n} - a_{d_n}(\boldsymbol{\lambda})}{\Delta\boldsymbol{\lambda}}. \quad (4.102)$$

Note that the difference quotient in Eq. (4.100), and hence the above expression for $\mathcal{A}_{d_n}(\boldsymbol{\lambda})$, only becomes exact for $\Delta\boldsymbol{\lambda} \rightarrow 0$. If we instead choose $\Delta\boldsymbol{\lambda} := \Delta\boldsymbol{\lambda}_j$ and plug the resulting approximation of $\mathcal{A}_{d_n}(\boldsymbol{\lambda})$ into Eq. (4.98), we get

$$\mathcal{W}_{d_n}(C) \approx i \sum_{j=1}^I (\mathbb{1}_{d_n} - a_{d_n}(\boldsymbol{\lambda}_j)), \quad (4.103)$$

where the $\Delta\boldsymbol{\lambda}_j$ cancel out and the elements of $a_{d_n}(\boldsymbol{\lambda}_j)$ take the simple form

$$a_{d_n,kl}(\boldsymbol{\lambda}_j) = \langle n_k(\boldsymbol{\lambda}_j) | n_l(\boldsymbol{\lambda}_j - \Delta\boldsymbol{\lambda}_j) \rangle = \langle n_k(\boldsymbol{\lambda}_j) | n_l(\boldsymbol{\lambda}_{j-1}) \rangle \quad (4.104)$$

due to $\boldsymbol{\lambda}_{j-1} = \boldsymbol{\lambda}_j - \Delta\boldsymbol{\lambda}_j$. With this, the WZ phase factor from Eq. (4.56) becomes

$$\mathcal{U}_{d_n}(C) = \mathcal{P} \exp [i \mathcal{W}_{d_n}(C)] \approx \mathcal{P} \exp \left[- \sum_{j=1}^I (\mathbb{1}_{d_n} - a_{d_n}(\boldsymbol{\lambda}_j)) \right]. \quad (4.105)$$

Even though the overlap matrices $a_{d_n}(\boldsymbol{\lambda}_j)$ do not generally commute, $[a_{d_n}(\boldsymbol{\lambda}_j), a_{d_n}(\boldsymbol{\lambda}_k)] \neq 0$ for $j \neq k$, when $d_n > 1$, the exponential in Eq. (4.105) can still be rewritten as

$$\mathcal{P} \exp \left[\sum_{j=1}^I (a_{d_n}(\boldsymbol{\lambda}_j) - \mathbb{1}_{d_n}) \right] = \mathcal{P} \prod_{j=1}^I \exp [a_{d_n}(\boldsymbol{\lambda}_j) - \mathbb{1}_{d_n}] \quad (4.106)$$

under the path ordering operator $\mathcal{P}$. As evident from the middle expression in Eq. (4.104), we then have

$$\lim_{\Delta\boldsymbol{\lambda}_j \rightarrow 0} a_{d_n}(\boldsymbol{\lambda}_j) \rightarrow \mathbb{1}_{d_n}, \quad (4.107)$$

so that the individual exponents on the right-hand side of Eq. (4.106) satisfy

$$\lim_{\Delta\boldsymbol{\lambda}_j \rightarrow 0} (a_{d_n}(\boldsymbol{\lambda}_j) - \mathbb{1}_{d_n}) \rightarrow 0. \quad (4.108)$$

For sufficiently high resolutions I and subsequently small $\Delta\boldsymbol{\lambda}_j$ we can therefore approximate the individual exponentials in Eq. (4.106) by

$$\exp [a_{d_n}(\boldsymbol{\lambda}_j) - \mathbb{1}_{d_n}] \approx \mathbb{1}_{d_n} + a_{d_n}(\boldsymbol{\lambda}_j) - \mathbb{1}_{d_n} = a_{d_n}(\boldsymbol{\lambda}_j) \quad (4.109)$$

and we get the formula

$$\mathcal{U}_{d_n}(C) \approx \mathcal{P} \prod_{j=1}^I a_{d_n}(\boldsymbol{\lambda}_j) \stackrel{(\diamond)}{\equiv} \prod_{j=1}^I a_{d_n}(\boldsymbol{\lambda}_{I+1-j}). \quad (4.110)$$

For $(\diamond)$, we explicitly carried out the path ordering

$$\mathcal{P} \prod_{j=1}^I a_{d_n}(\boldsymbol{\lambda}_j) = a_{d_n}(\boldsymbol{\lambda}_I) a_{d_n}(\boldsymbol{\lambda}_{I-1}) \cdots a_{d_n}(\boldsymbol{\lambda}_1) = \prod_{j=1}^I a_{d_n}(\boldsymbol{\lambda}_{I+1-j}). \quad (4.111)$$

In Chaps. 9 and 10 we use Eq. (4.110) to analyse the Berry and WZ phase factors of (non-)degenerate BdG vacua over the parameter manifold $\Lambda \equiv \mathbb{S}_\phi^1$ associated to the superconducting phase $\lambda \equiv \phi$. Specifically, Chap. 9 deals with the Abelian $d_0 = 1$ “Sombrero” Berry phases computed via the overlaps

$$a_1(\phi_j) = {}^P_b \langle 0(\phi_j) | 0(\phi_{j-1}) \rangle_b^P, \quad (4.112)$$

between instantaneous BdG vacua $|0(\phi_j)\rangle_b^P$ for different superconducting phases $\phi_j \in \mathbb{S}_\phi^1$. In Chap. 10, we extend this analysis to the non-Abelian $d_0 > 1$ WZ phase factors, which are in turn based on the overlap matrices

$$a_{d_0, mn}(\phi_j) = \langle 0_m(\phi_j) | 0_n(\phi_{j-1}) \rangle \equiv \langle m(\phi_j) | n(\phi_{j-1}) \rangle \quad (4.113)$$

between (almost) degenerate low-energy BdG Fock states $|n(\phi_j)\rangle$ where $n = 0, \dots, d_0 - 1$. Concretely, we discuss $d_0 = 2$ and $d_0 = 4$. In both cases, the BdG vacua and Fock states are constructed as detailed in Sec. 5.3 and the overlaps are computed using the Robledo Pfaffian formula reviewed in Sec. 5.4.

5 – Bogoliubov–de Gennes Theory

Bogoliubov–de Gennes (BdG) theory was developed to find solutions in Bardeen–Cooper–Schrieffer (BCS) theory, the first microscopic description of superconductivity. BCS theory famously explains superconductivity as a result of electron pairing due to an attractive electron-electron interaction at the Fermi surface. Since solving the interacting-electron problem is very challenging, BCS theory introduces a mean-field approximation to simplify it. The resulting mean-field BCS Hamiltonian can then be diagonalised using BdG theory, which introduces coherent superpositions of particles and holes known as Bogoliubov quasiparticles. In the BdG description, the superconducting ground state takes the form of a Bogoliubov quasiparticle vacuum, which, due to the fact that the quasiparticles combine elementary particle and hole excitations, is not an empty Fermi sea, but rather a macroscopic condensate of all states with even particle numbers.

In the following, we first give a brief introduction to BCS theory to provide some context for BdG theory as well as a reference for conventional and unconventional superconductivity. After that, we discuss the general formalism of BdG theory, describing Bogoliubov transformations, the notion of the Bogoliubov vacuum, and overlaps between Bogoliubov states. Finally, we discuss the possibility and significance of Majorana modes in BdG theories. Throughout this chapter we follow Refs. [82] and [83].

5.1 Bardeen–Cooper–Schrieffer Theory of Superconductivity

Between 1911 and 1957, superconductivity was an experimentally observed phenomenon lacking a proper theoretical description. This problem was partially solved when Bardeen, Cooper and Schrieffer proposed a microscopical model that turned out to be so successful it is still taught today [84]. This part is based on Ref. [82].

At its core, BCS theory shows that any attractive interaction between electrons near the Fermi surface of an electron gas can cause a ground state instability that energetically prioritises the formation of bound Cooper pairs of electrons with opposite momentum and spin. In conventional superconductors, this attractive interaction is a result of electron-phonon interactions. The usual start Hamiltonian reads

$$H = \sum_{\mathbf{k}, \alpha} \xi_{\mathbf{k}} c_{\mathbf{k}\alpha}^\dagger c_{\mathbf{k}\alpha} + \frac{1}{N} \sum_{\mathbf{k}, \mathbf{k}'} V_{\mathbf{k}\mathbf{k}'} c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}'\downarrow} c_{\mathbf{k}'\uparrow}, \quad (5.1)$$

where $c_{\mathbf{k}\alpha}^\dagger$ creates an electron with momentum $\mathbf{k}$ and spin $\alpha \in \{\uparrow, \downarrow\}$ and where $\xi_{\mathbf{k}} = \varepsilon_{\mathbf{k}} - \mu$ is the electronic energy dispersion $\varepsilon_{\mathbf{k}}$ shifted by a chemical potential μ . The second term is a crude model for the phonon-mediated attractive interaction causing the Cooper instability. It only takes into account scattering between time-reversed single-particle states $|\mathbf{k}, \uparrow\rangle$ and $|\mathbf{k}, \downarrow\rangle$ and is governed by an interaction strength

$$V_{\mathbf{k}\mathbf{k}'} = \begin{cases} -V & \text{for } |\xi_{\mathbf{k}}|, |\xi_{\mathbf{k}'}| \lesssim \omega_D \\ 0 & \text{else,} \end{cases} \quad (5.2)$$

which is constant for electrons with energies $|\xi_{\mathbf{k}}| \lesssim \omega_D$ relative to the Fermi energy and vanishes otherwise. Here, ω_D denotes the Debye frequency of the system.¹ The standard BCS Hamiltonian results when we apply the mean-field approximation

$$\langle c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}'\downarrow} c_{\mathbf{k}'\uparrow} \rangle \approx \langle c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger \rangle c_{-\mathbf{k}'\downarrow} c_{\mathbf{k}'\uparrow} + c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger \langle c_{-\mathbf{k}'\downarrow} c_{\mathbf{k}'\uparrow} \rangle - \langle c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger \rangle \langle c_{-\mathbf{k}'\downarrow} c_{\mathbf{k}'\uparrow} \rangle, \quad (5.3)$$

giving

$$H_{\text{BCS}} = \sum_{\mathbf{k}, \alpha} \xi_{\mathbf{k}} c_{\mathbf{k}\alpha}^\dagger c_{\mathbf{k}\alpha} - \sum_{\mathbf{k}} \left(\Delta_{\mathbf{k}} c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger + \text{H.c.} \right) + E_0, \quad (5.4)$$

¹The Debye frequency is the typical frequency of the phonon modes mediating the attractive electron-electron interaction in conventional superconductors.

where

$$\Delta_{\mathbf{k}} := -\frac{1}{N} \sum_{\mathbf{k}'} V_{\mathbf{k}\mathbf{k}'} \langle c_{-\mathbf{k}\downarrow} c_{\mathbf{k}\uparrow} \rangle \quad \text{and} \quad E_0 := -\frac{1}{N} \sum_{\mathbf{k}, \mathbf{k}'} V_{\mathbf{k}\mathbf{k}'} \langle c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger \rangle \langle c_{-\mathbf{k}\downarrow} c_{\mathbf{k}\uparrow} \rangle \quad (5.5)$$

have to be determined self-consistently using the BCS ground state that we will introduce shortly. The anomalous expectation values $\langle c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger \rangle$ and $\langle c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \rangle$ are non-zero only in the superconducting state. In that case, $\langle c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \rangle$ is interpreted as the momentum-spin representation of the Cooper pair created by $c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger$. As was mentioned in the introduction, H_{BCS} can be diagonalised using a Bogoliubov transformation. Anticipating a result presented later in this chapter, such a transformation takes the form

$$b_{\mathbf{k}\uparrow} := u_{\mathbf{k}} c_{\mathbf{k}\uparrow} + v_{\mathbf{k}} c_{-\mathbf{k}\downarrow}^\dagger \quad \text{and} \quad b_{-\mathbf{k}\downarrow} := u_{\mathbf{k}} c_{-\mathbf{k}\downarrow} - v_{\mathbf{k}} c_{\mathbf{k}\uparrow}^\dagger. \quad (5.6)$$

The complex coefficient functions $u_{\mathbf{k}}$ and $v_{\mathbf{k}}$ must satisfy $|u_{\mathbf{k}}|^2 + |v_{\mathbf{k}}|^2 = 1$ to ensure that the Bogoliubov quasiparticle operators in Eq. (5.6) obey fermionic anticommutator relations. The standard choice is

$$u_{\mathbf{k}} = \sqrt{\frac{1}{2} \left(1 + \frac{\xi_{\mathbf{k}}}{\sqrt{\xi_{\mathbf{k}}^2 + |\Delta_{\mathbf{k}}|^2}} \right)} \quad \text{and} \quad v_{\mathbf{k}} = \sqrt{\frac{1}{2} \left(1 - \frac{\xi_{\mathbf{k}}}{\sqrt{\xi_{\mathbf{k}}^2 + |\Delta_{\mathbf{k}}|^2}} \right)} e^{i\phi}, \quad (5.7)$$

where the complex phase factor $e^{i\phi}$ of $v_{\mathbf{k}}$ is given by the complex phase factor of the complex gap function in polar coordinates, i.e. $\Delta_{\mathbf{k}} = \Delta_0(\mathbf{k})e^{i\phi}$. We verify that Eq. (5.6) diagonalises H_{BCS} by plugging its inverse transformation

$$c_{\mathbf{k}\uparrow} = u_{\mathbf{k}}^* b_{\mathbf{k}\uparrow} + v_{\mathbf{k}} b_{-\mathbf{k}\downarrow}^\dagger \quad \text{and} \quad c_{-\mathbf{k}\downarrow}^\dagger = u_{\mathbf{k}} b_{-\mathbf{k}\downarrow}^\dagger - v_{\mathbf{k}}^* b_{\mathbf{k}\uparrow} \quad (5.8)$$

into Eq. (5.4), yielding

$$H_{\text{BCS}} = \sum_{\mathbf{k}, \alpha} E_{\mathbf{k}} b_{\mathbf{k}\alpha}^\dagger b_{\mathbf{k}\alpha} + \mathcal{E}_0 \quad (5.9)$$

with

$$E_{\mathbf{k}} = \sqrt{\xi_{\mathbf{k}}^2 + |\Delta_{\mathbf{k}}|^2} \quad \text{and} \quad \mathcal{E}_0 = \sum_{\mathbf{k}} (\xi_{\mathbf{k}} - E_{\mathbf{k}} + E_0). \quad (5.10)$$

This shows why $\Delta_{\mathbf{k}}$ is (usually) called the BCS gap function: it gaps the Bogoliubov quasiparticle spectrum even when the electronic energy dispersion vanishes. It also makes clear that the Bogoliubov quasiparticles themselves do *not* correspond to the Cooper pairs, but rather to excitations from the superconducting ground state. To see how this ground state emerges, we note that the Bogoliubov quasiparticle spectrum indicated in Eq. (5.10) is positive semi-definite. As a consequence, we can write the ground state of Eq. (5.9) as a Bogoliubov quasiparticle vacuum

$$|\text{GS}\rangle_{\text{BCS}} := |0\rangle_b, \quad (5.11)$$

in which no Bogoliubov quasiparticle is occupied. One way to determine $|0\rangle_b$ is as a product state

$$\begin{aligned} |0\rangle_b &= \prod_{\mathbf{k}} b_{\mathbf{k}\uparrow}^\dagger b_{-\mathbf{k}\downarrow} |0\rangle \\ &= \prod_{\mathbf{k}} \left(u_{\mathbf{k}} c_{\mathbf{k}\uparrow} + v_{\mathbf{k}} c_{-\mathbf{k}\downarrow}^\dagger \right) \left(u_{\mathbf{k}} c_{-\mathbf{k}\downarrow} - v_{\mathbf{k}} c_{\mathbf{k}\uparrow}^\dagger \right) |0\rangle \\ &= \prod_{\mathbf{k}} \left(-u_{\mathbf{k}} v_{\mathbf{k}} c_{\mathbf{k}\uparrow} c_{\mathbf{k}\uparrow}^\dagger + \cancel{u_{\mathbf{k}}^2 c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow}^\dagger} - v_{\mathbf{k}}^2 c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger + v_{\mathbf{k}} u_{\mathbf{k}} \cancel{c_{-\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}\downarrow}} \right) |0\rangle \\ &= \prod_{\mathbf{k}} \left(-u_{\mathbf{k}} v_{\mathbf{k}} (1 - \cancel{c_{\mathbf{k}\uparrow}^\dagger c_{\mathbf{k}\uparrow}}) - v_{\mathbf{k}}^2 c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) |0\rangle \\ &= \prod_{\mathbf{k}} (-v_{\mathbf{k}}) \prod_{\mathbf{k}} (u_{\mathbf{k}} + v_{\mathbf{k}} c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle, \end{aligned} \quad (5.12)$$

where $|0\rangle$ denotes the electronic vacuum and where $\mathcal{N} := \prod_{\mathbf{k}} (-v_{\mathbf{k}})$ determines the norm ${}_b\langle 0|0\rangle_b = \mathcal{N}^2$. A normalised BCS ground state is therefore given by

$$|\text{GS}\rangle_{\text{BCS}} = \prod_{\mathbf{k}} (u_{\mathbf{k}} + v_{\mathbf{k}} c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle, \quad (5.13)$$

which is a coherent superposition of all available even-parity states – a condensate of Cooper pairs.

So far, we have assumed that Cooper pairs are always formed by two electrons of opposite momentum and spin. As a result, the total spin S of the Cooper pair vanishes and it forms a spin singlet. Importantly, this assumption becomes questionable in the vicinity of ferromagnetic order, which favours parallel spin alignment and therefore triplet Cooper pairs with a total spin of $S = 1$. Such situations are very likely to exist in nature. One prominent example is the superfluidity of ^3He , which is generally attributed to triplet Cooper pairs of charge-neutral ^3He atoms, rather than electrons [82]. In order to describe these more exotic types of superconductivity, one has to generalise the BCS approach to superconductivity. This is done by allowing for arbitrary spin couplings, starting with

$$H = \sum_{\mathbf{k}, \alpha} \xi_{\mathbf{k}} c_{\mathbf{k}\alpha}^\dagger c_{\mathbf{k}\alpha} + \frac{1}{N} \sum_{\substack{\mathbf{k}, \mathbf{k}' \\ \alpha, \beta, \gamma, \delta}} V_{\mathbf{k}\mathbf{k}'\alpha\beta\gamma\delta} c_{\mathbf{k}\alpha}^\dagger c_{-\mathbf{k}\beta}^\dagger c_{-\mathbf{k}'\gamma} c_{\mathbf{k}'\delta}, \quad (5.14)$$

in which the only generic symmetries of $V_{\mathbf{k}\mathbf{k}'\alpha\beta\gamma\delta}$ are imposed by the fermionic anticommutation relations:

$$V_{\mathbf{k}\mathbf{k}'\alpha\beta\gamma\delta} = -V_{\mathbf{k}-\mathbf{k}'\alpha\beta\delta\gamma} = -V_{-\mathbf{k}\mathbf{k}'\beta\alpha\gamma\delta} = V_{-\mathbf{k}-\mathbf{k}'\beta\alpha\delta\gamma}. \quad (5.15)$$

In particular, Eq. (5.14) does not exclude Cooper pairs of electrons with the same spin. In a weak-coupling approach with an interaction that is attractive, i.e. $V_{\mathbf{k}\mathbf{k}'\alpha\beta\gamma\delta} < 0$ in an energy range $E_F \pm \varepsilon_c$ defined by a finite cutoff-energy $|\varepsilon_c| \ll |E_F|$, we can repeat the above construction and get a mean-field Hamiltonian

$$H_{\text{BCS}} = \sum_{\mathbf{k}, \alpha} \xi_{\mathbf{k}} c_{\mathbf{k}\alpha}^\dagger c_{\mathbf{k}\alpha} - \frac{1}{2} \sum_{\mathbf{k}, \alpha, \beta} \left[\Delta_{\mathbf{k}\alpha\beta} c_{\mathbf{k}\alpha}^\dagger c_{-\mathbf{k}\beta}^\dagger + \Delta_{\mathbf{k}\alpha\beta}^* c_{\mathbf{k}\alpha} c_{-\mathbf{k}\beta} \right] + E_0, \quad (5.16)$$

where the generalised gap equation

$$\Delta_{\mathbf{k}\alpha\beta} = -\frac{1}{N} \sum_{\mathbf{k}', \gamma, \delta} V_{\mathbf{k}\mathbf{k}'\alpha\beta\gamma\delta} \langle c_{-\mathbf{k}'\gamma} c_{\mathbf{k}'\delta} \rangle \quad (5.17)$$

and the offset energy

$$E_0 = -\frac{1}{N} \sum_{\substack{\mathbf{k}, \mathbf{k}' \\ \alpha, \beta, \gamma, \delta}} V_{\mathbf{k}\mathbf{k}'\alpha\beta\gamma\delta} \langle c_{\mathbf{k}\alpha}^\dagger c_{-\mathbf{k}\beta}^\dagger \rangle \langle c_{-\mathbf{k}'\gamma} c_{\mathbf{k}'\delta} \rangle \quad (5.18)$$

have to be determined self-consistently. In order to analyse the type of superconductivity that a given pairing mechanism $V_{\mathbf{k}\mathbf{k}'\alpha\beta\gamma\delta}$ produces, we characterise the structure of the resulting Cooper pairs based on their expectation values $\langle c_{-\mathbf{k}\alpha} c_{\mathbf{k}\beta} \rangle$. As mentioned earlier, the $\langle c_{-\mathbf{k}\alpha} c_{\mathbf{k}\beta} \rangle$ are usually interpreted as the momentum-spin representation of the Cooper pair wave functions, motivating the notation

$$\psi_{\mathbf{k}, \alpha\beta} \equiv \langle c_{-\mathbf{k}\alpha} c_{\mathbf{k}\beta} \rangle. \quad (5.19)$$

By definition, $\psi_{\mathbf{k}, \alpha\beta}$ is odd under an exchange of electrons within the Cooper pair, i.e.

$$\psi_{\mathbf{k}, \alpha\beta} = -\psi_{-\mathbf{k}, \beta\alpha}, \quad (5.20)$$

and one can use this antisymmetry to analyse the symmetry properties of $\psi_{\mathbf{k}, \alpha\beta}$ under the individual exchange of momentum- and spin-quantum numbers in Eq. (5.19). Specifically, we can formally separate the Cooper pair wave function as

$$\psi_{\mathbf{k}, \alpha\beta} \equiv \phi_{\mathbf{k}} \cdot \chi_{\alpha\beta} \quad (5.21)$$

into an orbital part $\phi_{\mathbf{k}}$ and a spin part $\chi_{\alpha\beta}$. If the orbital angular momentum L of the Cooper pair wave function is a good quantum number, their orbital parity is given by $(-1)^L$ and we can write

$$\begin{aligned} \phi_{\mathbf{k}} = \phi_{-\mathbf{k}} &\iff \chi_{\alpha\beta} = -\chi_{\beta\alpha} && \text{giving} && L = 0, 2, 4, \dots && \text{and} && S = 0 \\ \phi_{\mathbf{k}} = -\phi_{-\mathbf{k}} &\iff \chi_{\alpha\beta} = \chi_{\beta\alpha} && \text{giving} && L = 1, 3, 5, \dots && \text{and} && S = 1, \end{aligned} \quad (5.22)$$

where $S = 0$ means that the Cooper pairs form (antisymmetric) spin singlets and $S = 1$ means that the Cooper pairs form (symmetric) spin triplets. Generally, superconductivity with $L = 0$ Cooper pairs is

referred to as conventional superconductivity, while superconductivity with $L > 0$ Cooper pairs is called unconventional. In combination with the fundamental symmetries of the interaction given in Eq. (5.15), the Cooper pair symmetries in Eq. (5.22) fully determine the symmetry properties of the BCS gap function from Eq. (5.17). We get

$$\Delta_{\mathbf{k}\alpha\beta} \stackrel{(5.15)}{=} -\Delta_{-\mathbf{k}\beta\alpha} \stackrel{(5.22)}{=} \begin{cases} \Delta_{-\mathbf{k}\alpha\beta} = -\Delta_{\mathbf{k}\beta\alpha} & \text{for } L = 0, 2, 4, \dots \text{ and } S = 0 \\ -\Delta_{-\mathbf{k}\alpha\beta} = \Delta_{\mathbf{k}\beta\alpha} & \text{for } L = 1, 3, 5, \dots \text{ and } S = 1. \end{cases} \quad (5.23)$$

Generally, the gap function can be written as a complex 2×2 matrix

$$\Delta_{\mathbf{k}} \equiv \begin{pmatrix} \Delta_{\mathbf{k}\uparrow\uparrow} & \Delta_{\mathbf{k}\uparrow\downarrow} \\ \Delta_{\mathbf{k}\downarrow\uparrow} & \Delta_{\mathbf{k}\downarrow\downarrow} \end{pmatrix} \quad (5.24)$$

in the spin indices. We can find generic parametrisations of $\Delta_{\mathbf{k}}$ using the symmetry relations in Eq. (5.23). For even L and $S = 0$, we only need one function $f(\mathbf{k})$ satisfying $f(\mathbf{k}) = f(-\mathbf{k})$ to write

$$\Delta_{\mathbf{k}}^{S=0} \equiv \begin{pmatrix} 0 & f(\mathbf{k}) \\ -f(\mathbf{k}) & 0 \end{pmatrix} = if(\mathbf{k})\sigma_y, \quad (5.25)$$

where σ_y denotes the y -Pauli matrix. For odd L and $S = 1$ we need three functions $d_1(\mathbf{k}), d_2(\mathbf{k}), d_3(\mathbf{k})$ satisfying $d_i(\mathbf{k}) = -d_i(-\mathbf{k})$ each. If we combine these as $\mathbf{d}(\mathbf{k}) = (d_1(\mathbf{k}), d_2(\mathbf{k}), d_3(\mathbf{k}))^\top$ we can write

$$\Delta_{\mathbf{k}}^{S=1} \equiv \begin{pmatrix} -d_1(\mathbf{k}) + id_2(\mathbf{k}) & d_3(\mathbf{k}) \\ d_3(\mathbf{k}) & d_1(\mathbf{k}) + id_2(\mathbf{k}) \end{pmatrix} = i(\mathbf{d}(\mathbf{k})\boldsymbol{\sigma})\sigma_y, \quad (5.26)$$

where $\boldsymbol{\sigma}$ denotes the vector of Pauli matrices. The most general BCS gap matrix is therefore given by

$$\Delta_{\mathbf{k}} = \Delta_{\mathbf{k}}^{S=0} + \Delta_{\mathbf{k}}^{S=1} = i(f(\mathbf{k})\mathbb{1}_2 + \mathbf{d}(\mathbf{k})\boldsymbol{\sigma})\sigma_y \quad (5.27)$$

with $f(\mathbf{k})$ and $\mathbf{d}(\mathbf{k})$ as before. However, in most cases, either singlet or triplet pairing dominates this expression and it is valid to speak of singlet or triplet superconductivity.

The self-consistent gap equation of the conventional BCS Hamiltonian with phonon-mediated attractive interaction given in Eq. (5.4) turns out $\mathbf{k}$ -independent, so that conventional superconductivity corresponds to $L = 0$ and $S = 0$ Cooper pairs [82, 85]. Borrowing from the naming conventions of atomic orbitals, this type of singlet superconductivity is also known as s -wave superconductivity. In fact, most known superconductors are singlet superconductors, including all conventional s -wave superconductors, cuprates (d -wave) and the pnictides ($s_{\pm}$ -wave) [86].² In comparison, spin-triplet superconductors are relatively rare. The strongest candidate is the aforementioned ^3He . Other candidates include UPt_3 and Sr_2RuO_4 , which are also prominent contenders for a special type of superconductivity known as chiral superconductivity [86]. Chiral superconductors are characterised by a complex superconducting gap function $\Delta_{\mathbf{k}}$, whose phase winds around some axis on the Fermi surface. The direction of this winding gives these systems a definite handedness, or chirality, for which they are named. The simplest example is a so-called $k_x + ik_y$ gap function, whose phase winds up by $\pm 2\pi$ as $\mathbf{k}$ follows a closed path around the k_z -axis. As loosely suggested by the presence of a characteristic winding number, chiral superconductivity is a topological state of matter, supporting topological zero modes at certain defects. For instance, there are boundary modes, which disperse across the superconducting gap for $\mathbf{k}$ parallel to the boundary. Another example are non-dispersive vortex-bound zero modes, which may exist in some triplet chiral superconductors. These arise at cores of superfluid vortices, topological defects where the phase of the order parameter winds by 2π in real space, and can be understood as Majorana quasiparticles, i.e. excitations that are their own antiparticles. The same applies to the dispersive boundary mode for the specific $\mathbf{k}$ at which it passes through zero energy. Even more striking than their formal Majorana property is the fact that the vortex-bound Majorana zero modes obey non-Abelian Ising anyon statistics. This makes triplet chiral superconductors a promising platform for topological quantum computing, as we will discuss in greater detail later on.

²An $s_{\pm}$ -wave gap has an s -wave symmetry, but changes its sign as a function of $\mathbf{k}$ [82].

5.2 The Bogoliubov–de Gennes Formalism

The Bogoliubov–de Gennes (BdG) formalism introduces a special class of unitary transformations called Bogoliubov transformations. These are characterised by their compatibility with a symmetry-like particle-hole structure that distinguishes a certain class of quadratic fermionic Hamiltonians known as BdG Hamiltonians. A BdG Hamiltonian is usually written as

$$H_{\text{BdG}} = \frac{1}{2} \Psi^\dagger h_{\text{BdG}} \Psi, \quad (5.28)$$

where $\Psi = (c_1 \dots c_d c_1^\dagger \dots c_d^\dagger)^\top =: (\mathbf{c} \mathbf{c}^\dagger)^\top$ is called a Nambu spinor, and where

$$h_{\text{BdG}} = \begin{pmatrix} T & \Delta^\dagger \\ \Delta & -T^* \end{pmatrix} \quad (5.29)$$

denotes the $2d \times 2d$ BdG matrix, whose blocks are the $d \times d$ single-particle hopping matrix T and the $d \times d$ anomalous gap matrix Δ . In terms of the unitary PHC operator Ξ from Sec. 3, the particle-hole structure of H_{BdG} and h_{BdG} is given by

$$\Xi H_{\text{BdG}} \Xi = -H_{\text{BdG}}^* \quad \text{and} \quad \Xi h_{\text{BdG}} \Xi = -h_{\text{BdG}}^*, \quad (5.30)$$

cf. Eq. (3.38). A Bogoliubov transformation is then a unitary $2d \times 2d$ matrix

$$B = \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix}, \quad (5.31)$$

which fulfils the particle-hole compatibility condition

$$\Xi B \Xi = B^* \quad (5.32)$$

and diagonalises h_{BdG} as

$$B^\dagger h_{\text{BdG}} B = E_{\text{BdG}}, \quad (5.33)$$

where $E_{\text{BdG}} = \sigma_z \otimes E$ with $E = \text{diag}(E_1, \dots, E_d)$ and $E_i \geq 0$. Thus, it provides a transformation

$$H_{\text{BdG}} = \frac{1}{2} \Psi^\dagger h_{\text{BdG}} \Psi = \frac{1}{2} \Psi^\dagger B E_{\text{BdG}} B^\dagger \Psi =: \frac{1}{2} \Phi^\dagger E_{\text{BdG}} \Phi = \sum_{j=1}^d E_j b_j^\dagger b_j \quad (5.34)$$

of the BdG Hamiltonian H_{BdG} . Here, we defined Bogoliubov-transformed Nambu spinors

$$\Phi =: \begin{pmatrix} \mathbf{b} \\ \mathbf{b}^\dagger \end{pmatrix} = \begin{pmatrix} U^\dagger & V^\dagger \\ V^\top & U^\top \end{pmatrix} \begin{pmatrix} \mathbf{c} \\ \mathbf{c}^\dagger \end{pmatrix} = B^\dagger \Psi, \quad (5.35)$$

which give rise to Bogoliubov quasiparticle creation and annihilation operators $b_j^\dagger$ and b_j . The main practical features of Bogoliubov transformations are twofold. The most important one is that the resulting Bogoliubov quasiparticle operators obey the usual fermionic anticommutation relations

$$\{b_j, b_k^\dagger\} = \delta_{jk} \quad \text{and} \quad \{b_j, b_k\} = \{b_j^\dagger, b_k^\dagger\} = 0. \quad (5.36)$$

The second noteworthy feature is that the quasiparticles associated to $b_j^\dagger$ and the quasiholes associated to b_j are naturally related through particle-hole conjugation and have opposite energies $+E_j$ and $-E_j$. The fact that $E_j \geq 0$ in Eq. (5.34) allows us to write the many-body BdG ground state as a vacuum

$$|\text{GS}\rangle_{\text{BdG}} = |0\rangle_b \quad (5.37)$$

of the Bogoliubov quasiparticles. One way to obtain this vacuum is as a simple *product state*

$$|0\rangle_b = \prod_{j=1}^d b_j |0\rangle \equiv |0\rangle_b^{\text{p}}, \quad (5.38)$$

where $|0\rangle$ denotes the true electronic vacuum, which serves as a canonic reference state.

In the following, we will discuss the BdG formalism in more detail. We start by addressing the setting of Nambu space and the shortcomings of naive unitary diagonalisation procedures therein. Next, we analyse how particle-hole conjugation Ξ can be used to resolve these issues, giving rise to Bogoliubov transformations. Finally, we provide a construction guide for these transformations, including definitions of the Bogoliubov vacuum and a formula for many-body overlaps between Bogoliubov states.

Given an d -dimensional single-particle Hilbert space $\mathcal{H}$ spanned by d single particle states $|\psi_1\rangle, \dots, |\psi_d\rangle$, we will understand the associated $2d$ -dimensional Nambu space $\mathbf{H}$ as the generalised Hilbert space

$$\mathbf{H} := \mathcal{H} \oplus \mathcal{H}^*, \quad (5.39)$$

spanned by both the single-particle states $|\psi_1\rangle, \dots, |\psi_d\rangle \in \mathcal{H}$ and their duals $\langle\psi_1|, \dots, \langle\psi_d| \in \mathcal{H}^*$, the “single-hole” states. In second quantisation, this corresponds to a basis

$$\mathcal{B} = \{c_1, \dots, c_d, c_1^\dagger, \dots, c_d^\dagger\} \quad (5.40)$$

of all elementary annihilation and creation operators, respectively. We will stick to the language of creation and annihilation operators from now on. A BdG Hamiltonian is then a quadratic Hermitian operator

$$H_{\text{BdG}} = \sum_{i,j} h_{ij} d_i d_j \quad (5.41)$$

on $\mathbf{H}$, where $d_i, d_j \in \mathcal{B}$. This formulation naturally incorporates anomalous terms $\propto c_i^\dagger c_j^\dagger$ and $\propto c_i c_j$ that are present in mean-field Hamiltonians like the BCS Hamiltonian we discussed earlier. In particular, the BdG formulation in Nambu space only adds value if such anomalous terms are present. Note that we can specify Eq. (5.41) as

$$H_{\text{BdG}} = \sum_{i,j} \left(A_{ij} c_i^\dagger c_j + B_{ij} c_i c_j^\dagger + C_{ij} c_i c_j + D_{ij} c_i^\dagger c_j^\dagger \right), \quad (5.42)$$

where we introduced four auxiliary $N \times N$ matrices A, B, C, D governing the four types of quadratic terms that are possible in Eq. (5.41). The only restrictions on A, B, C, D come from the Hermiticity of H_{BdG} , which requires

$$A = A^\dagger, \quad B = B^\dagger, \quad C = -D^* \quad (5.43)$$

and the fermionic anticommutation relations of the creation and annihilation operators, which require

$$A = -B^\top, \quad C = -C^\top, \quad D = -D^\top \quad (5.44)$$

for a consistent and non-trivial definition of H_{BdG} . Combining Eqs. (5.43) and (5.44), we get conditions

$$B = -A^* \quad \text{and} \quad D = C^\dagger. \quad (5.45)$$

If we substitute Eq. (5.45) into Eq. (5.42) and rename $A = T$ and $C = \Delta$ we get

$$\begin{aligned} H_{\text{BdG}} &= \sum_{i,j} \left(T_{ij} c_i^\dagger c_j - T_{ij}^* c_i c_j^\dagger + \Delta_{ij} c_i c_j - \Delta_{ij}^* c_i^\dagger c_j^\dagger \right) \\ &= \sum_{i,j} \left(T_{ij} c_i^\dagger c_j - T_{ij}^* c_i c_j^\dagger + \Delta_{ij} c_i c_j + \Delta_{ij}^\dagger c_i^\dagger c_j^\dagger \right), \end{aligned} \quad (5.46)$$

which readily aligns with the usual notation from mean-field BCS theory, where T represents the single-particle hopping matrix and Δ denotes anomalous gap matrix. As mentioned in the beginning, a defining feature of these Hamiltonians is their symmetry-like particle-hole structure. To formulate this, we use the unitary PHC operator Ξ from before. It swaps the single-particle creation and annihilation operators as

$$\Xi c_j \Xi^{-1} = c_j^\dagger \quad \text{and} \quad \Xi c_j^\dagger \Xi^{-1} = c_j. \quad (5.47)$$

Note that Ξ defines an involution, since repeating the transformation twice is equivalent to the identity transformation. This shows that Ξ is self-inverse, meaning $\Xi^{-1} = \Xi$. For this reason, we will simply write Ξ for both the PHC transformation and its inverse from now on. The transformation of H_{BdG} under Ξ is then given by

$$\begin{aligned}\Xi H_{\text{BdG}} \Xi &= \Xi \sum_{i,j} \left(T_{ij} c_i^\dagger c_j - T_{ij}^* c_i c_j^\dagger + \Delta_{ij} c_i c_j - \Delta_{ij}^* c_i^\dagger c_j^\dagger \right) \Xi \\ &= \sum_{i,j} \left(T_{ij} \Xi c_i^\dagger \Xi \Xi c_j \Xi - T_{ij}^* \Xi c_i \Xi \Xi c_j^\dagger \Xi + \Delta_{ij} \Xi c_i \Xi \Xi c_j \Xi - \Delta_{ij}^* \Xi c_i^\dagger \Xi \Xi c_j^\dagger \Xi \right) \\ &= \sum_{i,j} \left(T_{ij} c_i c_j^\dagger - T_{ij}^* c_i^\dagger c_j + \Delta_{ij} c_i^\dagger c_j^\dagger - \Delta_{ij}^* c_i c_j \right) \\ &= -H_{\text{BdG}}^*,\end{aligned}\tag{5.48}$$

where we defined the complex conjugate Hamiltonian H_{BdG}^* as the original Hamiltonian H_{BdG} with complex conjugate coefficients. The particle-hole relation

$$\Xi H_{\text{BdG}} \Xi = -H_{\text{BdG}}^* \tag{5.49}$$

is often referred to as a symmetry in the literature. We avoid this terminology for a number of reasons. Most notably, Eq. (5.49) relates H_{BdG} to its complex conjugate H_{BdG}^* rather than to H_{BdG} itself. This makes the interpretation of Ξ as a quantum symmetry operation very difficult. One way to fix this is to resort to the antiunitary PHC operator $\bar{\Xi}$ from Eq. (3.36), which turns Eq. (5.49) into a relation

$$\bar{\Xi} H_{\text{BdG}} \bar{\Xi} = -H_{\text{BdG}} \tag{5.50}$$

between H_{BdG} and itself. However, even if we accept the now unusual antilinearity of $\bar{\Xi}$, we are still left with an unconventional symmetry because Eq. (5.50) tells us that $\bar{\Xi}$ anticommutes with H_{BdG} , whereas typical symmetry operators commute with it. Finally, both Eqs. (5.49) and (5.50) are a direct result of the Hermiticity of H_{BdG} and the fermionic algebra of the c -operators and thus not related to an additional symmetry of the system in the traditional sense. For this reason, Eq. (5.50) is sometimes called a “tautological” particle-hole symmetry as well. For a an in depth discussion of this see [78]. Even though the interpretation of the particle-hole structure given in Eqs. (5.49) and (5.50) is a bit subtle, it is far from useless. For instance, Eq. (5.50) tells us that if $|\psi\rangle$ is an eigenstate of H_{BdG} with energy E , then its PHC counterpart $|\phi\rangle = \bar{\Xi} |\psi\rangle$ is also an eigenstate of H_{BdG} but with energy $-E$. This can be seen by

$$\begin{aligned}H_{\text{BdG}} |\phi\rangle &= H_{\text{BdG}} \bar{\Xi} |\psi\rangle \\ &\stackrel{(\diamond)}{=} -\bar{\Xi} H_{\text{BdG}} |\psi\rangle \\ &\stackrel{(\star)}{=} -\bar{\Xi} E |\psi\rangle \\ &= -E \bar{\Xi} |\psi\rangle \\ &= -E |\phi\rangle,\end{aligned}\tag{5.51}$$

where we first utilised Eq. (5.49) in $(\diamond)$ and then $H_{\text{BdG}} |\psi\rangle = E |\psi\rangle$ in $(\star)$. The result of the particle-hole structure of H_{BdG} is therefore not a degeneracy of eigenstates, but a particle-hole pairing between eigenstates of opposite energies. One advantage of the BdG setting is that it allows us to discuss H_{BdG} and Ξ in terms of their matrix representations. Specifically, the Hamiltonian matrix h_{BdG} is defined via

$$H_{\text{BdG}} = \Psi^\dagger h_{\text{BdG}} \Psi = \begin{pmatrix} c^\dagger & c \end{pmatrix} \begin{pmatrix} T & \Delta^\dagger \\ \Delta & -T^* \end{pmatrix} \begin{pmatrix} c \\ c^\dagger \end{pmatrix}, \tag{5.52}$$

where h_{BdG} can be extracted from Eq. (5.46), while the matrix representation of Ξ reads

$$\Xi = \sigma_x \otimes \mathbb{1}_d = \begin{pmatrix} 0 & \mathbb{1}_d \\ \mathbb{1}_d & 0 \end{pmatrix}, \tag{5.53}$$

which is given in the standard Nambu basis Eq. (5.40). Here, $\mathbb{1}_d$ denotes the $d \times d$ identity matrix. A quick sanity check shows that

$$\Xi h_{\text{BdG}} \Xi = \begin{pmatrix} 0 & \mathbb{1}_d \\ \mathbb{1}_d & 0 \end{pmatrix} \begin{pmatrix} T & \Delta^* \\ \Delta & -T^* \end{pmatrix} \begin{pmatrix} 0 & \mathbb{1}_d \\ \mathbb{1}_d & 0 \end{pmatrix} = \begin{pmatrix} -T^* & \Delta \\ -\Delta^* & T \end{pmatrix} = -h_{\text{BdG}}^*, \quad (5.54)$$

where we have used $\Delta^\top = -\Delta$ to write $\Delta^\dagger = -\Delta^*$ to make the relation more obvious. Note also that the definition of H_{BdG} that we gave in Eq. (5.28) has an overall prefactor of one half compared to Eq. (5.52). This prefactor is a matter of “bookkeeping” – it represents a correction of the Nambu redundancy, which arises from the fact that every physical degree of freedom is taken into account twice, once as a particle and once as a hole. It appears naturally when one starts with a quadratic mean-field Hamiltonian and brings it into Nambu space.

The powerful thing about Eq. (5.52) is that it allows us to diagonalise H_{BdG} by diagonalising its matrix representation h_{BdG} . Thus, we look for a unitary $2d \times 2d$ matrix U , transforming

$$U^\dagger h_{\text{BdG}} U = \text{diag}(E_1, \dots, E_{2d}) =: E_{\text{BdG}}, \quad (5.55)$$

where the columns of U correspond to the eigenvectors of h_{BdG} as usual. At this point, the Nambu redundancy, which allowed us to translate the diagonalisation procedure of H_{BdG} into a linear algebra problem, causes some subtleties that have to be taken care of. In a regular fermionic single-particle Hilbert space, every unitary diagonalisation of a particle-number conserving Hamiltonian H guarantees that the resulting eigenstates form another basis of the Hilbert space. Specifically, we can write H as

$$H = \sum_{i,j} T_{ij} c_i^\dagger c_j = (c_1^\dagger \dots c_d^\dagger) T \begin{pmatrix} c_1 \\ \vdots \\ c_d \end{pmatrix} =: \psi^\dagger T \psi \quad (5.56)$$

and find unitary matrices U that diagonalise T as

$$U^\dagger T U = \text{diag}(E_1, \dots, E_d) =: E_T \quad (5.57)$$

to get a diagonalisation

$$H = \psi^\dagger T \psi = \psi^\dagger U E_T U^\dagger \psi = \phi^\dagger E_T \phi = \sum_j E_j \lambda_j^\dagger \lambda_j \quad (5.58)$$

of H . The unitarity of U is now necessary and sufficient to ensure that the new fermionic creation and annihilation operators

$$\lambda_i := U_{ij}^\dagger c_j \quad \text{and} \quad \lambda_i^\dagger := U_{ji} c_j^\dagger, \quad (5.59)$$

defined via Eq. (5.58), satisfy the usual fermionic anticommutation relations

$$\{\lambda_j, \lambda_k^\dagger\} = \delta_{jk} \quad \text{and} \quad \{\lambda_j, \lambda_k\} = \{\lambda_j^\dagger, \lambda_k^\dagger\} = 0. \quad (5.60)$$

This property is lost for BdG Hamiltonians on Nambu space because unitary diagonalisation matrices are generally going to mix particle and hole states. To see this, we write a generic unitary transformation of the Nambu basis as

$$\lambda_{(i,x)} := \sum_{j=1}^d \left(U_{(i,x)(j,1)}^\dagger c_j + U_{(i,x)(j,2)}^\dagger c_j^\dagger \right), \quad (5.61)$$

where we introduced multi-indices (i, x) with $i = 1, \dots, d$ and $x = 1, 2$ that reflect the Nambu redundancy. Concretely, $(1, 1), \dots, (d, 1)$ refer to the single-*particle* part of the Nambu basis, while $(1, 2), \dots, (d, 2)$ refer to its single-*hole* part. Even though such a unitary transformation produces an orthonormal basis of $2d$ energy eigenstates in Nambu space, the underlying physical theory can only accomodate d distinct

eigenstates. We will come back to this later. For now we convince ourselves that unitarity still partly ensures fermionicity of the $\lambda_{(i,x)}$ since

$$\begin{aligned}
\{\lambda_{(i,x)}, \lambda_{(j,y)}^\dagger\} &= \sum_{n,k=1}^d \left\{ \left(U_{(i,x)(n,1)}^\dagger c_n + U_{(i,x)(n,2)}^\dagger c_n^\dagger \right), \left(U_{(j,y)(k,1)}^\dagger c_k + U_{(j,y)(k,2)}^\dagger c_k^\dagger \right) \right\} \\
&= \sum_{n,k=1}^d \left(U_{(i,x)(n,1)}^\dagger U_{(j,y)(k,1)}^\dagger \{c_n, c_k^\dagger\} + U_{(i,x)(n,1)}^\dagger U_{(j,y)(k,2)}^\dagger \{c_n, c_k\} \right. \\
&\quad \left. + U_{(i,x)(n,2)}^\dagger U_{(j,y)(k,1)}^\dagger \{c_n^\dagger, c_k^\dagger\} + U_{(i,x)(n,2)}^\dagger U_{(j,y)(k,2)}^\dagger \{c_n^\dagger, c_k\} \right) \\
&= \sum_{n,k=1}^d \left(U_{(i,x)(n,1)}^\dagger U_{(j,y)(k,1)}^\dagger \delta_{nk} + U_{(i,x)(n,2)}^\dagger U_{(j,y)(k,2)}^\dagger \delta_{nk} \right) \\
&= \sum_{n=1}^d \left(U_{(i,x)(n,1)}^\dagger U_{(j,y)(n,1)}^\dagger + U_{(i,x)(n,2)}^\dagger U_{(j,y)(n,2)}^\dagger \right) \\
&= \sum_{n=1}^d \left(U_{(i,x)(n,1)}^\dagger U_{(n,1)(j,y)} + U_{(i,x)(n,2)}^\dagger U_{(n,2)(j,y)} \right) \\
&= \sum_{(n,z)} U_{(i,x)(n,z)}^\dagger U_{(n,z)(j,y)} \\
&= (U^\dagger U)_{(i,x)(j,y)} \\
&= \delta_{(i,x)(j,y)} .
\end{aligned} \tag{5.62}$$

This tells us that the unitarity of U continues to guarantee the correct anticommutation relations between the new annihilation and creation operators. However, the anticommutators *among* the new annihilation operators (and analogously among the new creation operators) are given by

$$\begin{aligned}
\{\lambda_{(i,x)}, \lambda_{(j,y)}\} &= \sum_{n,k=1}^d \left\{ \left(U_{(i,x)(n,1)}^\dagger c_n + U_{(i,x)(n,2)}^\dagger c_n^\dagger \right), \left(U_{(j,y)(k,1)}^\dagger c_k + U_{(j,y)(k,2)}^\dagger c_k^\dagger \right) \right\} \\
&= \sum_{n,k=1}^d \left(U_{(i,x)(n,1)}^\dagger U_{(j,y)(k,1)}^\dagger \{c_n, c_k\} + U_{(i,x)(n,1)}^\dagger U_{(j,y)(k,2)}^\dagger \{c_n, c_k^\dagger\} \right. \\
&\quad \left. + U_{(i,x)(n,2)}^\dagger U_{(j,y)(k,1)}^\dagger \{c_n^\dagger, c_k\} + U_{(i,x)(n,2)}^\dagger U_{(j,y)(k,2)}^\dagger \{c_n^\dagger, c_k^\dagger\} \right) \\
&= \sum_{n,k=1}^d \left(U_{(i,x)(n,1)}^\dagger U_{(j,y)(k,2)}^\dagger \delta_{nk} + U_{(i,x)(n,2)}^\dagger U_{(j,y)(k,1)}^\dagger \delta_{nk} \right) \\
&= \sum_{n=1}^d \left(U_{(i,x)(n,1)}^\dagger U_{(j,y)(n,2)}^\dagger + U_{(i,x)(n,2)}^\dagger U_{(j,y)(n,1)}^\dagger \right) \\
&= \sum_{(n,z)} U_{(i,x)(n,z)}^\dagger U_{(j,y)(n,\bar{z})}^\dagger .
\end{aligned} \tag{5.63}$$

Here, we introduced a shorthand notation $\bar{z}$ in the last line that refers to the opposite particle-hole index of z , i.e. $\bar{1} = 2$ and $\bar{2} = 1$, for a more compact notation. The problem is, that the unitarity of U is not enough to ensure that these vanish. Of course, salvation lies in the only other structure at hand, namely the particle-conjugation structure of the BdG Hamiltonian. It allows us to transform Eq. (5.55) as

$$\begin{aligned}
&\Xi U^\dagger h_{\text{BdG}} U \Xi = \Xi E_{\text{BdG}} \Xi \\
\iff &\Xi U^\dagger \Xi \Xi h_{\text{BdG}} \Xi \Xi U \Xi = \Xi E_{\text{BdG}} \Xi \\
\iff &-V^\dagger h_{\text{BdG}}^* V = \tilde{E}_{\text{BdG}} ,
\end{aligned} \tag{5.64}$$

where we used $\Xi^2 = \mathbb{1}_{2d}$, plugged in the particle-hole conjugation of h_{BdG} from Eq. (5.54), and finally defined

$$V := \Xi U \Xi \quad \text{and} \quad \tilde{E}_{\text{BdG}} := \Xi E_{\text{BdG}} \Xi . \tag{5.65}$$

The particle-hole conjugation matrix Ξ transforms any $2d \times 2d$ matrix M as

$$\Xi M \Xi = \begin{pmatrix} 0 & \mathbb{1}_d \\ \mathbb{1}_d & 0 \end{pmatrix} \begin{pmatrix} A & B \\ C & D \end{pmatrix} \begin{pmatrix} 0 & \mathbb{1}_d \\ \mathbb{1}_d & 0 \end{pmatrix} = \begin{pmatrix} D & C \\ B & A \end{pmatrix}, \quad (5.66)$$

i.e. it swaps the $d \times d$ diagonal and off-diagonal blocks of M . It follows that $\tilde{E}_{\text{BdG}}$ is again a diagonal matrix that contains the eigenvalues of h_{BdG} , only in a different order. Moreover, Eq. (5.51) ensures that the same holds for $-\tilde{E}_{\text{BdG}}$. Thus, there exists a real orthogonal permutation matrix R such that

$$-\tilde{E}_{\text{BdG}} = R^\dagger E_{\text{BdG}} R, \quad (5.67)$$

where we wrote $R^\dagger = R^\top = R^{-1}$ to make the upcoming expression a bit simpler. Taking the complex conjugate of the final line of Eq. (5.64) therefore yields the relation

$$V^{*\dagger} h_{\text{BdG}} V^* = -\tilde{E}_{\text{BdG}} = R^\dagger E_{\text{BdG}} R \stackrel{(\diamond)}{=} R^\dagger U^\dagger h_{\text{BdG}} U R, \quad (5.68)$$

where we used that $\tilde{E}_{\text{BdG}}^* = \tilde{E}_{\text{BdG}}$ is a real matrix and where we plugged in Eq. (5.55) in $(\diamond)$. The definition Eq. (5.65) of V then yields the relation

$$UR = V^* = \Xi U^* \Xi \quad (5.69)$$

between the unitary diagonalisation matrix U and its complex conjugate U^* . We can use this relation to further analyse the problematic anticommutators between annihilation operators from Eq. (5.63). To do this, we first note that the multi-index notation allows us to write the matrix elements of Ξ as

$$\Xi_{(i,x)(j,y)} = \delta_{(i,\bar{x})(i,y)} = \delta_{(i,x)(i,\bar{y})} = \delta_{ij}(1 - \delta_{xy}), \quad (5.70)$$

such that

$$\begin{aligned} (\Xi A \Xi)_{(i,x)(j,y)} &= \sum_{(l,m)} \sum_{(p,q)} \Xi_{(i,x)(l,m)} A_{(l,m)(p,q)} \Xi_{(p,q)(j,y)} \\ &= \sum_{(l,m)} \sum_{(p,q)} \delta_{(i,\bar{x})(l,m)} A_{(l,m)(p,q)} \delta_{(p,q)(j,\bar{y})} \\ &= A_{(i,\bar{x})(j,\bar{y})} \end{aligned} \quad (5.71)$$

for any $2d \times 2d$ matrix A . With this, we can use Eq. (5.69) to relate the matrix elements of U and U^* as

$$U_{(i,x)(j,y)}^* = (\Xi U R \Xi)_{(i,x)(j,y)} = (UR)_{(i,\bar{x})(j,\bar{y})}. \quad (5.72)$$

If we plug this into Eq. (5.63) we obtain

$$\begin{aligned} \{\lambda_{(i,x)}, \lambda_{(j,y)}\} &= \sum_{(n,z)} U_{(i,x)(n,z)}^\dagger U_{(j,y)(n,\bar{z})}^\dagger \\ &= \sum_{(n,z)} U_{(i,x)(n,z)}^\dagger U_{(n,\bar{z})(j,y)}^* \\ &= \sum_{(n,z)} U_{(i,x)(n,z)}^\dagger (UR)_{(n,z)(j,\bar{y})} \\ &= \sum_{(n,z)} \sum_{(p,q)} U_{(i,x)(n,z)}^\dagger U_{(n,z)(p,q)} R_{(p,q)(j,\bar{y})} \\ &= \sum_{(p,q)} (U^\dagger U)_{(i,x)(p,q)} R_{(p,q)(j,\bar{y})} \\ &= \sum_{(p,q)} \delta_{(i,x)(p,q)} R_{(p,q)(j,\bar{y})} \\ &= R_{(i,x)(j,\bar{y})}. \end{aligned} \quad (5.73)$$

Taking the adjoint of this equation then tells us that

$$\{\lambda_{(i,x)}, \lambda_{(j,y)}\} = \{\lambda_{(i,x)}^\dagger, \lambda_{(j,y)}^\dagger\} = R_{(i,x)(j,\bar{y})} . \quad (5.74)$$

Since $R = 0$ is not an allowed permutation matrix, there is no way to obtain the usual fermionic anti-commutators

$$\{\lambda_{(i,x)}, \lambda_{(j,y)}\} = \{\lambda_{(i,x)}^\dagger, \lambda_{(j,y)}^\dagger\} = 0 . \quad (5.75)$$

However, this is expected due to the Nambu redundancy. In fact, even the original fermions only fulfil

$$\{d_{(i,x)}, d_{(j,y)}\} = \{d_{(i,x)}^\dagger, d_{(j,y)}^\dagger\} = \delta_{(i,x)(j,\bar{y})} \quad (5.76)$$

for $d_{i,1} = c_i$ and $d_{i,2} = c_i^\dagger$, where the only anomalous anticommutators

$$\{d_{(i,x)}, d_{(i,\bar{x})}\} = \{d_{(i,x)}^\dagger, d_{(i,\bar{x})}^\dagger\} = 1 \quad (5.77)$$

are the ones between the particle and hole representatives of the same fermion, e.g.

$$\{d_{(i,1)}, d_{(i,\bar{1})}\} = \{d_{(i,1)}, d_{(i,2)}\} = \{c_i, c_i^\dagger\} = 1 . \quad (5.78)$$

In order to realise the anticommutation relations from Eq. (5.76) for the $\lambda_{(i,x)}$ and $\lambda_{(i,x)}^\dagger$ operators, the permutation matrix R has to be precisely the identity matrix

$$R_{(i,x)(j,y)} = \delta_{(i,x)(j,y)} , \quad (5.79)$$

since then Eq. (5.74) becomes

$$\{\lambda_{(i,x)}, \lambda_{(j,y)}\} = \{\lambda_{(i,x)}^\dagger, \lambda_{(j,y)}^\dagger\} = \delta_{(i,x)(j,\bar{y})} , \quad (5.80)$$

which tells us that the $\lambda_{(i,x)}$ are ordinary fermions as long as we restrict their indices to either Nambu sector, i.e. to the particle sector ($x = 1$) or the hole sector ($x = 2$). Strictly speaking we could run sensible physics with any collection $C = \{\lambda_{(i_1,x_1)}, \dots, \lambda_{(i_d,x_d)}\}$ of d particle-hole independent fermion operators, where particle-hole independence refers to the property that C cannot contain both the particle- and the hole-version of one and the same fermion. We conclude that a diagonalisation transformation yields proper quasi-fermions if and only if its unitary matrix representation U satisfies a much stricter version of Eq. (5.69), namely

$$U \stackrel{!}{=} \Xi U^* \Xi . \quad (5.81)$$

We define Bogoliubov transformations as the class of unitary transformations that satisfy this condition. This leaves us with the problem of how to *construct* a Bogoliubov transformation for a given BdG Hamiltonian. One possible starting point reveals itself upon closer inspection of Eq. (5.81). It provides us with a strikingly simple relation between U and U^* . In particular, it almost looks like $U = U^*$, which would make U a real unitary and therefore orthogonal matrix. Even though Eq. (5.81) does not quite make U itself an orthogonal matrix, it can still be used to establish a unitary equivalence between U and an orthogonal matrix, identifying the Bogoliubov transformations as the orthogonal subgroup $O(2d)$ of the unitary group $U(2d)$. Specifically, one can find a unitary matrix $\mathcal{U}$ with which the particle-hole conjugation matrix Ξ can be decomposed as

$$\Xi = \mathcal{U} \mathcal{U}^\top . \quad (5.82)$$

If we plug this into Eq. (5.81) we get

$$U \stackrel{!}{=} \mathcal{U} \mathcal{U}^\top U^* \mathcal{U} \mathcal{U}^\dagger \iff \mathcal{U}^\dagger U \mathcal{U} \stackrel{!}{=} \mathcal{U}^\top U^* \mathcal{U}^* = (\mathcal{U}^\dagger U \mathcal{U})^* \iff O \stackrel{!}{=} O^* , \quad (5.83)$$

where we defined

$$O = \mathcal{U}^\dagger U \mathcal{U} . \quad (5.84)$$

The basic idea is then the following. Maybe we can find the desired Bogoliubov diagonalisation

$$U^\dagger h_{\text{BdG}} U = E_{\text{BdG}} \quad (5.85)$$

of a given BdG matrix h_{BdG} by determining the orthogonal transformation O that solves the $\mathcal{U}$ -related problem

$$O^\top a_{\text{BdG}} O = C_{\text{BdG}} , \quad (5.86)$$

which results from Eq. (5.85) via

$$\mathcal{U}^\dagger U^\dagger \mathcal{U} \mathcal{U}^\dagger h_{\text{BdG}} \mathcal{U} \mathcal{U}^\dagger U \mathcal{U} = \mathcal{U}^\dagger E_{\text{BdG}} \mathcal{U} \quad (5.87)$$

upon defining

$$a_{\text{BdG}} = \mathcal{U}^\dagger h_{\text{BdG}} \mathcal{U} \quad \text{and} \quad C_{\text{BdG}} = \mathcal{U}^\dagger E_{\text{BdG}} \mathcal{U} . \quad (5.88)$$

This may sound like a bold hope, but it turns out to be quite a workable solution. The reason is that the matrix a_{BdG} comes out imaginary, skew-symmetric and Hermitian, and the orthogonal similarity transformation O that brings the associated real, skew-symmetric and Hermitian matrix

$$a'_{\text{BdG}} := i a_{\text{BdG}} \quad (5.89)$$

into the required form C_{BdG} can be found by established algorithms. Let us go through this procedure in more detail. For a start, we show that the unitary matrix $\mathcal{U}$ from Eq. (5.82) exists. The simple form $\Xi = \sigma_x \otimes \mathbb{1}_d$ of the particle-hole conjugation matrix Ξ suggests the ansatz

$$\mathcal{U} = \sum_{i \in I} a_i \sigma_i \otimes \mathbb{1}_d , \quad (5.90)$$

where $a_i \in \mathbb{C}$ and where we introduced the index set $I = \{x, y, z\}$ for a more convenient notation. The unitarity of $\mathcal{U}$ requires

$$\begin{aligned} \mathcal{U} \mathcal{U}^\dagger &= \sum_{i,j \in I} a_i a_j^* \sigma_i \sigma_j^\dagger \otimes \mathbb{1}_d, \\ &= \sum_{i,j \in I} a_i a_j^* \sigma_i \sigma_j \otimes \mathbb{1}_d, \\ &\stackrel{(\diamond)}{=} \sum_{i,j \in I} a_i a_j^* (\delta_{ij} \mathbb{1}_2 + i \epsilon_{ijk} \sigma_k) \otimes \mathbb{1}_d, \\ &\stackrel{!}{=} \mathbb{1}_{2d} , \end{aligned} \quad (5.91)$$

where we have used the product identity $\sigma_i \sigma_j = (\delta_{ij} \mathbb{1}_2 + i \epsilon_{ijk} \sigma_k)$ of the Pauli matrices in $(\diamond)$.³ This yields the simple constraints

$$1 \stackrel{!}{=} \sum_{i \in I} |a_i|^2 \quad \text{and} \quad 0 \stackrel{!}{=} \sum_{i,j \in I} a_i a_j^* \epsilon_{ijk} \quad (5.92)$$

for all $k = x, y, z$. The latter is equivalent to $a_i a_j^* \stackrel{!}{=} a_i^* a_j$, which can be further simplified to $a_i a_j^* \stackrel{!}{\in} i\mathbb{R}$. Similarly, we get

$$\begin{aligned} \mathcal{U} \mathcal{U}^\top &= \sum_{i,j \in I} a_i a_j \sigma_i \sigma_j^\top \otimes \mathbb{1}_d \\ &= \sum_{i,j \in I} a_i a_j (-1)^{\delta_{ij}} (\delta_{ij} \mathbb{1}_2 + i \epsilon_{ijk} \sigma_k) \otimes \mathbb{1}_d \\ &\stackrel{!}{=} \sigma_x \otimes \mathbb{1}_d , \end{aligned} \quad (5.93)$$

³Here, the Levi-Civita symbol ϵ_{ijk} refers to the standard order $ijk = xyz$, i.e. we have $\epsilon_{xyz} = \epsilon_{yzx} = \epsilon_{zxy} = 1$ and $\epsilon_{yxz} = \epsilon_{xzy} = \epsilon_{zyx} = -1$.

where we wrote

$$\sigma_j^\Gamma = (-1)^{\delta_{jy}} \sigma_j \quad (5.94)$$

to capture the symmetry (skew-symmetry) of the x and z (the y) Pauli matrix. From this, we obtain the condition

$$\sigma_x \stackrel{!}{=} \sum_{i,j \in I} a_i a_j (-1)^{\delta_{jy}} (\delta_{ij} \mathbb{1}_2 + i \epsilon_{ijk} \sigma_k), \quad (5.95)$$

which translates to

$$0 \stackrel{!}{=} \sum_{i \in I} a_i^2 (-1)^{\delta_{iy}} \quad \text{and} \quad i \sum_{i \in I} a_i a_j (-1)^{\delta_{jy}} \epsilon_{ijk} \stackrel{!}{=} \begin{cases} 1 & k = x \\ 0 & \text{else} \end{cases} \quad (5.96)$$

Once more, we are only interested in a particular solution for this. We can choose $a_x = 0$, which immediately ensures that the only non-zero term in the right-hand condition above is the $k = x$ term that we are interested in. With $a_x = 0$, the left-hand condition becomes $0 \stackrel{!}{=} a_y^2 - a_z^2$, so we need $a_y = a_z$ to satisfy it. In order to determine the particular value of a_y and a_z , we look at the $k = x$ condition

$$1 \stackrel{!}{=} i(a_y a_z + a_z a_y) \quad \implies \quad a_y a_z \stackrel{!}{=} -\frac{i}{2}, \quad (5.97)$$

which is readily satisfied by

$$a_y = a_z = \frac{e^{-i\pi/4}}{\sqrt{2}}. \quad (5.98)$$

A quick check confirms that these coefficient choices are in line with Eq. (5.92), so the result is going to be unitary. Combined we get the solution

$$\mathcal{U} = \frac{e^{-i\pi/4}}{\sqrt{2}} (\sigma_y + \sigma_z) \otimes \mathbb{1}_d = \frac{e^{-i\pi/4}}{\sqrt{2}} \begin{pmatrix} \mathbb{1}_d & -i\mathbb{1}_d \\ i\mathbb{1}_d & -\mathbb{1}_d \end{pmatrix}. \quad (5.99)$$

Thus, we can always decompose the particle-hole conjugation matrix Ξ as given in in Eq. (5.82), allowing us to identify the subgroup of Bogoliubov transformations as the orthogonal subgroup $O(2d)$ of the unitary group $U(2d)$ via the isomorphism

$$\begin{aligned} \mathcal{O} : U(2d) &\rightarrow O(2d) < U(2d) \\ U &\mapsto \mathcal{U}^\dagger U \mathcal{U}, \end{aligned} \quad (5.100)$$

defined as the conjugation with $\mathcal{U}$. Note that this identification tells us that all Bogoliubov transformations B satisfy $\det(B) = \pm 1$, allowing a distinction between proper ($\det(B) = +1$) and improper ($\det(B) = -1$) Bogoliubov transformations. With Eq. (5.99) at hand we can determine the aforementioned auxiliary matrix

$$a_{\text{BdG}} = \mathcal{U}^\dagger h_{\text{BdG}} \mathcal{U} \quad (5.101)$$

and consider its properties. Namely, we find that it fulfils

$$\begin{aligned} a_{\text{BdG}}^* &= (\mathcal{U}^\dagger h_{\text{BdG}} \mathcal{U})^* \\ &= \mathcal{U}^\Gamma h_{\text{BdG}}^* \mathcal{U}^* \\ &= \mathcal{U}^\Gamma (-\Xi h_{\text{BdG}} \Xi) \mathcal{U}^* \\ &= -\mathcal{U}^\Gamma (\mathcal{U}^* \mathcal{U}^\dagger) h_{\text{BdG}} (\mathcal{U} \mathcal{U}^\Gamma) \mathcal{U}^* \\ &= -(\mathcal{U}^\Gamma \mathcal{U}^*) \mathcal{U}^\dagger h_{\text{BdG}} \mathcal{U} (\mathcal{U}^\Gamma \mathcal{U}^*) \\ &= -\mathcal{U}^\dagger h_{\text{BdG}} \mathcal{U} \\ &= -a_{\text{BdG}}, \end{aligned} \quad (5.102)$$

where we used Eq. (5.54) in the third line, the defining relation $\mathcal{U}\mathcal{U}^\top = \Xi = \Xi^\dagger = \mathcal{U}^*\mathcal{U}^\dagger$ of $\mathcal{U}$ in the fourth line, and the usual unitarity $\mathcal{U}^\dagger\mathcal{U} = \mathbb{1}_{2d}$ giving $\mathcal{U}^\top\mathcal{U}^* = \mathbb{1}_{2d}^* = \mathbb{1}_{2d}$ in the sixth line. An analogous calculation shows that

$$\begin{aligned} a_{\text{BdG}}^\top &= (\mathcal{U}^\dagger h_{\text{BdG}} \mathcal{U})^\top \\ &= \mathcal{U}^\top h_{\text{BdG}}^\top \mathcal{U}^* \\ &= \mathcal{U}^\top h_{\text{BdG}}^* \mathcal{U}^* \\ &= -a_{\text{BdG}} , \end{aligned} \tag{5.103}$$

where we used the Hermiticity of h_{BdG} to rewrite $h_{\text{BdG}}^\top = h_{\text{BdG}}^*$, allowing us to reuse the previous calculation. Thus, the conjugation of the BdG matrix h_{BdG} by $\mathcal{U}$ yields an imaginary, skew-symmetric and Hermitian matrix a_{BdG} . As mentioned in the beginning, we can use a_{BdG} to define a real and skew-symmetric matrix

$$a'_{\text{BdG}} := ia_{\text{BdG}} . \tag{5.104}$$

Real skew-symmetric $2d \times 2d$ matrices like a'_{BdG} are unitarily diagonalised and possess an imaginary point spectrum of pairwise conjugate eigenvalues

$$s_p(a'_{\text{BdG}}) = \{\pm i\lambda_1, \dots, \pm i\lambda_d\} . \tag{5.105}$$

Furthermore, it is a well-known fact that there exists an orthogonal transformation

$$O^\top a'_{\text{BdG}} O = N , \tag{5.106}$$

which brings a'_{BdG} into the canonical form

$$N = \begin{pmatrix} 0 & \lambda_1 & \dots & 0 & 0 \\ -\lambda_1 & 0 & \dots & 0 & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \dots & 0 & \lambda_d \\ 0 & 0 & \dots & -\lambda_d & 0 \end{pmatrix} = \Lambda \otimes (i\sigma_y) , \tag{5.107}$$

where we defined the diagonal matrix

$$\Lambda := \text{diag}(\lambda_1, \dots, \lambda_d) \tag{5.108}$$

of non-negative⁴ eigenvalues of a'_{BdG} . Note that this form of N exploits that a'_{BdG} is always even-dimensional. As a consequence, any zero eigenvalues appear also in pairs and can be incorporated into the 2×2 block-diagonal structure in the form of 2×2 zero blocks. Furthermore, the tensor product structure of Eq. (5.107) looks a bit unfamiliar because the first factor is d -dimensional and the second factor is two-dimensional. This is precisely the opposite order of the tensor structure of Ξ and $\mathcal{U}$, where the first factor is two-dimensional and the second factor is d -dimensional. Indeed, the above form of N corresponds to a different choice of Nambu basis, namely

$$\tilde{\mathcal{B}} = \{c_1, c_1^\dagger, \dots, c_d, c_d^\dagger\} . \tag{5.109}$$

Nambu spinors $\tilde{\Psi}$ given in this basis are related to Nambu spinors Ψ given in the canonical basis from Eq. (5.40) via

$$\tilde{\Psi} = Y \Psi , \tag{5.110}$$

where Y is another simple real orthogonal permutation matrix. We are going to use this transformation to change between Nambu bases whenever the discussion in basis $\tilde{\mathcal{B}}$ is beneficial, for example when

⁴Of course, the eigenvalues of a_{BdG} and $a'_{\text{BdG}} = ia_{\text{BdG}}$ only differ by a factor of i . Also, technically this definition of Λ is only correct if there are no zero eigenvalues. If there are, they are going to come in even numbers and Λ will only contain one half of them.

the standard algorithms from linear algebra provide us with a matrix of the form Eq. (5.107). With Eq. (5.110) we can then turn Eq. (5.106) into

$$Y^\top O^\top a'_{\text{BdG}} O Y = Y^\top N Y \quad \implies \quad W^\top a'_{\text{BdG}} W = M, \quad (5.111)$$

where we defined $W := OY$ and $M := Y^\top N Y$, which is given by

$$M = (i\sigma_y) \otimes \Lambda. \quad (5.112)$$

Now M matches our canonical particle-hole tensor product structure. If we multiply the right-hand expression in Eq. (5.111) by $\mathcal{U}$ from the left and by $\mathcal{U}^\dagger$ from the right we get

$$\begin{aligned} \mathcal{U} M \mathcal{U}^\dagger &= \mathcal{U} W^\top a'_{\text{BdG}} W \mathcal{U}^\dagger \\ &= \mathcal{U} W^\top (\mathcal{U}^\dagger \mathcal{U}) (i a_{\text{BdG}}) (\mathcal{U}^\dagger \mathcal{U}) W \mathcal{U}^\dagger \\ &= i (\mathcal{U} W^\top \mathcal{U}^\dagger) (\mathcal{U} a_{\text{BdG}} \mathcal{U}^\dagger) (\mathcal{U} W \mathcal{U}^\dagger) \\ &= i (\mathcal{U} W \mathcal{U}^\dagger)^\dagger (\mathcal{U} a_{\text{BdG}} \mathcal{U}^\dagger) (\mathcal{U} W \mathcal{U}^\dagger), \end{aligned} \quad (5.113)$$

which, using Eq. (5.101), gives

$$(\mathcal{U} W \mathcal{U}^\dagger)^\dagger h_{\text{BdG}} (\mathcal{U} W \mathcal{U}^\dagger) = -i \mathcal{U} M \mathcal{U}^\dagger, \quad (5.114)$$

and hence

$$B^\dagger h_{\text{BdG}} B = D, \quad (5.115)$$

where we defined

$$B = \mathcal{U} W \mathcal{U}^\dagger \quad \text{and} \quad D = -i \mathcal{U} M \mathcal{U}^\dagger. \quad (5.116)$$

We can determine the matrix D by plugging in the definition Eq. (5.99) of $\mathcal{U}$ as

$$\begin{aligned} D &= -i \left(\frac{e^{-i\pi/4}}{\sqrt{2}} (\sigma_y + \sigma_z) \otimes \mathbb{1}_d \right) ((i\sigma_y) \otimes \Lambda) \left(\frac{e^{i\pi/4}}{\sqrt{2}} (\sigma_y + \sigma_z) \otimes \mathbb{1}_d \right) \\ &= \frac{1}{2} ((\sigma_y + \sigma_z) \otimes \mathbb{1}_d) (\sigma_y \otimes \Lambda) ((\sigma_y + \sigma_z) \otimes \mathbb{1}_d) \\ &= \frac{1}{2} ((\sigma_y + \sigma_z) \sigma_y (\sigma_y + \sigma_z)) \otimes \Lambda \\ &= \frac{1}{2} ((\sigma_y + \sigma_z) (\mathbb{1}_2 + i\sigma_x)) \otimes \Lambda \\ &= \sigma_z \otimes \Lambda, \end{aligned} \quad (5.117)$$

which is a diagonal matrix

$$D = \begin{pmatrix} \Lambda & 0 \\ 0 & -\Lambda \end{pmatrix} \equiv E_{\text{BdG}} \quad (5.118)$$

with the eigenvalues of h_{BdG} on its diagonal. In particular, Λ was constructed such that it held only eigenvalues $E_i \geq 0$ so the two blocks of D readily group the positive and negative eigenenergies that we assigned to particle and hole states in the beginning. Thus, it is valid to write

$$B^\dagger h_{\text{BdG}} B = E_{\text{BdG}}, \quad (5.119)$$

highlighting that B is indeed a Bogoliubov transformation that diagonalises h_{BdG} as desired. The problem of finding a Bogoliubov transformation B that diagonalises a given BdG Hamiltonian h_{BdG} is therefore reduced to finding an orthogonal transformation O that brings $a'_{\text{BdG}} = i\mathcal{U}^\dagger h_{\text{BdG}} \mathcal{U}$ into the canonical form given in Eq. (5.107).

One reliable way to achieve this is by means of a fundamental result from linear algebra known as Schur decomposition. The Schur decomposition is a matrix decomposition algorithm that allows one to write any real (complex) square matrix M as

$$M = QSQ^{-1}, \quad (5.120)$$

where S is a real (complex) upper triangular matrix called the Schur form of M and where Q is an orthogonal (unitary) transformation matrix. The Schur decomposition is of great methodical and computational value because upper triangular matrices are generally much easier to handle than full square matrices. In case of the complex Schur decomposition, the matrix S has another very neat feature: its diagonal is made up of its eigenvalues, which, since S and M are similar, are also the eigenvalues of M . Here, we are going to be concerned with the real Schur decomposition, which is a little bit more subtle. The reason for this is that a general real square matrix M may have both real and complex eigenvalues, although the latter must appear as conjugate pairs. If a real square matrix M does have at least one conjugate pair of complex eigenvalues, there is no way to reduce it to an upper triangular form by means of a real (orthogonal) similarity transformation. Thus, the real Schur decomposition generally produces a matrix S , which is not quite in an upper triangular form. Instead, S has an upper block-triangular form

$$S = \begin{pmatrix} K_1 & \times & \cdots & \times \\ 0 & K_2 & \cdots & \times \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & K_d \end{pmatrix}, \quad (5.121)$$

where the K_j making up the diagonal are either real 1×1 matrices

$$K_j = m_j \quad (5.122)$$

containing a real eigenvalue m_j of M , or real 2×2 matrices of the form

$$K_j = \begin{pmatrix} a_j & b_j \\ -b_j & a_j \end{pmatrix}, \quad (5.123)$$

representing a conjugate pair $a_j \pm ib_j$ of complex eigenvalues of M . A proof of the real Schur decomposition theorem is included in App. A.2.

The utility of the Schur decomposition lies in its generality: any real (or complex) square matrix has a Schur form, that can be used to reduce the computational effort in numerical settings or improve efficiency in linear algebra problems. For example, the spectral theorem for real symmetric matrices, which states that every real symmetric matrix $A = A^\top$ is diagonalisable by an orthogonal transformation, may be regarded as a direct consequence of the Schur decomposition: the Schur decomposition theorem states that A is similar, through an orthogonal transformation, to an upper triangular block matrix S . Since S is similar to A , it has to have the same symmetries as A . In particular, $S^\top \stackrel{!}{=} S$, such that the only non-zero elements in S are the elements of the diagonal blocks. Since real symmetric matrices only have real eigenvalues, all diagonal blocks of S are 1×1 blocks of eigenvalues, such that A is diagonalised by the orthogonal transformation. Similarly, the Schur form S of a real, skew-symmetric matrix $A = -A^\top$ must also be real and skew-symmetric. As a result, the diagonal of S is bound to be zero and the only possibly non-zero elements are occupying the off-diagonals of its 2×2 blocks, i.e. all 1×1 blocks are equal to zero and the 2×2 blocks of the form Eq. (5.123) simplify to

$$K_j = \begin{pmatrix} 0 & k_j \\ -k_j & 0 \end{pmatrix}. \quad (5.124)$$

Note, that this immediately allows us to read off a few things we already know about the spectra of real, skew-symmetric matrices. Namely, the only real eigenvalues they can have are equal to zero, and their only non-zero eigenvalues are conjugate pairs of imaginary numbers, as can be seen from the eigenvalues $\pm ik_j$ of K_j in Eq. (5.124). Furthermore, the fact that the imaginary eigenvalues come in pairs tells us that all odd-dimensional real, skew-symmetric matrices must have an odd number⁵ of zero eigenvalues,

⁵In particular, at least one.

making them necessarily singular. Fortunately, the real skew-symmetric matrix a'_{BdG} that we obtain from a BdG matrix h_{BdG} is even by construction, so its Schur form is precisely a matrix

$$S = \begin{pmatrix} 0 & \lambda_1 & \cdots & 0 & 0 \\ -\lambda_1 & 0 & \cdots & 0 & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \cdots & 0 & \lambda_d \\ 0 & 0 & \cdots & -\lambda_d & 0 \end{pmatrix} = \Lambda \otimes (i\sigma_y), \quad (5.125)$$

where $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_d)$ is a diagonal matrix with the imaginary parts of independent eigenvalues of A on its diagonal. Since a'_{BdG} is even-dimensional, the number of zero eigenvalues has to be even as well, i.e. they too come in pairs and we can essentially treat them in the same way we treat the pairs of imaginary eigenvalues, allowing us to write S as in Eq. (5.125).

We conclude that the real Schur decomposition can be used to obtain the orthogonal similarity transformation O and the canonic form N , as defined in Eq. (5.106). Thus, the Schur decomposition provides us with everything we need to determine a Bogoliubov transformation B and spectral matrix E_{BdG} for a given BdG Hamiltonian. This is particularly useful for numerical purposes, as it enables us to build on well-established and stable implementations of the real Schur decomposition algorithm.

There is one last subtlety regarding the numerical output of a Schur decomposition algorithm, that we need to consider: a numerical implementation of the Schur decomposition may produce a Schur matrix S lacking the uniform structure we implied in Eq. (5.125). There are two reasons for this. The first is that zero eigenvalues appear as 1×1 blocks that may and generally will be arbitrarily positioned on the diagonal of S . In particular, zero blocks are not necessarily grouped together as in Eq. (5.125). The second is that there generally is no uniform sign convention for the 2×2 blocks of S . This means that S may feature 2×2 blocks of both the forms

$$K_j = \begin{pmatrix} 0 & k_j \\ -k_j & 0 \end{pmatrix} \quad \text{and} \quad K_j = \begin{pmatrix} 0 & -k_j \\ k_j & 0 \end{pmatrix}, \quad (5.126)$$

simultaneously. Both of these irregularities translate to an ambiguity in the definition of the auxiliary matrix Λ , cf. e.g. Eq. (5.125), holding half of the eigenvalues of our BdG matrix h_{BdG} . We have mentioned before, that we want to choose Λ such that it is positive-semidefinite because then the eigenvalue matrix E_{BdG} from Eq. (5.118) readily manifests the understanding of $(U, V)^\top$ as the positive energy Bogoliubov eigenvectors in B , allowing us to understand the many-body ground state in terms of a Bogoliubov vacuum state. Thus, we have to process the numerical Schur form accordingly. Say we get a numerical Schur decomposition

$$Q^\top a'_{\text{BdG}} Q = S \quad (5.127)$$

of our auxiliary BdG matrix. Then we have to determine an appropriate orthogonal transformation Z , such that

$$N := Z^\top S Z = |\Lambda| \otimes (i\sigma_y), \quad (5.128)$$

where

$$|\Lambda| = \text{diag}(|\lambda_1|, \dots, |\lambda_d|) \quad \text{with} \quad 0 \leq |\lambda_1| \leq \dots \leq |\lambda_d| \quad (5.129)$$

is the diagonal matrix of sorted individual eigenvalues that we wished for in Eq. (5.108). The block structure Eq. (5.125) of S tells us that

$$Z = \begin{pmatrix} F_1 & 0 & 0 & \cdots & 0 \\ 0 & F_2 & 0 & \cdots & 0 \\ 0 & 0 & F_3 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & F_d \end{pmatrix} \quad \text{with} \quad F_i = \begin{cases} \mathbb{1}_2 & \text{if } \lambda_i > 0 \\ \sigma_x & \text{if } \lambda_i < 0, \end{cases} \quad (5.130)$$

i.e. that Z is precisely the orthogonal transformation that flips the columns and rows of a 2×2 block if the corresponding superdiagonal element of S is negative, and does nothing otherwise. If we transform Eq. (5.127) using Z we get

$$Z^\top Q^\top a'_{\text{BdG}} Q Z = Z^\top S Z, \quad (5.131)$$

which eventually gives

$$O^\top a'_{\text{BdG}} O = N, \quad (5.132)$$

where we defined

$$O = QZ \quad \text{and} \quad N = |\Lambda| \otimes (i\sigma_y) = \begin{pmatrix} 0 & |\lambda_1| & \cdots & 0 & 0 \\ -|\lambda_1| & 0 & \cdots & 0 & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \cdots & 0 & |\lambda_d| \\ 0 & 0 & \cdots & -|\lambda_d| & 0 \end{pmatrix}, \quad (5.133)$$

which is precisely the form we looked for in Eq. (5.107). From here, we can obtain the desired Bogoliubov transformation as outlined above.

5.3 The Bogoliubov Vacuum

Let H_{BdG} be a BdG Hamiltonian and let B be a Bogoliubov transformation diagonalising it as

$$H_{\text{BdG}} = \sum_j E_j b_j^\dagger b_j, \quad (5.134)$$

where $E_j \geq 0$ and where the Bogoliubov quasiparticle operators $b_j^\dagger$ and b_j are defined via

$$\Phi = \begin{pmatrix} \mathbf{b} \\ \mathbf{b}^\dagger \end{pmatrix} = \begin{pmatrix} U^\dagger & V^\dagger \\ V^\top & U^\top \end{pmatrix} \begin{pmatrix} \mathbf{c} \\ \mathbf{c}^\dagger \end{pmatrix} = B^\dagger \Psi. \quad (5.135)$$

Here, Ψ (Φ) denotes the Nambu spinor of elementary (Bogoliubov) fermion operators defined in terms of the column vectors $\mathbf{c}$ ($\mathbf{b}$) of elementary (Bogoliubov) annihilation operators and the corresponding column vectors $\mathbf{c}^\dagger$ ($\mathbf{b}^\dagger$) of elementary (Bogoliubov) creation operators. Due to $E_j \geq 0$ in Eq. (5.134), we can express the many-body ground state of H_{BdG} as a vacuum state of Bogoliubov quasiparticles, i.e.

$$|\text{GS}\rangle_{\text{BdG}} = |0\rangle_b. \quad (5.136)$$

The Bogoliubov vacuum $|0\rangle_b$ is naturally defined via

$$b_j |0\rangle_b \stackrel{!}{=} 0 \quad (5.137)$$

for all annihilation operators b_j of Bogoliubov quasiparticles. So how do we come by such a Bogoliubov vacuum state? We are going to outline two approaches to defining a Bogoliubov vacuum. The first one stands out for its formal elegance. It is rooted in the representation theory of Lie groups and provides an insightful yet technical perspective. The second one covers simplified expressions that are well-known for their utility and established across many fields of application. The second approach will be the focus of this work. Before turning to it, however, it is instructive to first consider a more formal perspective.

From a group theoretical point of view, a Bogoliubov vacuum state $|0\rangle_b$ can be obtained using the representation theory of the Lie group $\text{O}(2d)$. We have previously shown that the group of unitary Bogoliubov transformations acting on a Nambu space $\mathbf{H} = \mathcal{H} \oplus \mathcal{H}^*$ of an d -dimensional complex Hilbert space $\mathcal{H}$ is isomorphic to the orthogonal group $\text{O}(2d)$ in $2d$ real dimensions, providing a partition of the Bogoliubov group into proper ($\det(B) = 1$) and improper ($\det(B) = -1$) Bogoliubov transformations.

The exponential map from Lie group theory allows us to express every proper Bogoliubov transformation $B \in \text{SO}(2d) < \text{O}(2d)$ as

$$B = \exp [X] , \quad (5.138)$$

where $X \in \mathfrak{o}(2d)$ is a suitable skew-symmetric real $2d \times 2d$ matrix from the Lie algebra $\mathfrak{o}(2d)$ of $\text{SO}(2d)$. For a given B , the matrix X is formally determined⁶ by the matrix logarithm, i.e. $X = \log(B)$. The Bogoliubov vacuum we are looking for is a distinguished state in the fermionic BdG Fock space. Generally, the Fock space $\mathcal{F}(\mathcal{H})$ of a single-particle Hilbert space $\mathcal{H}$ is formed by the action of polynomials in the single-particle creation operators $c_j^\dagger$ on a state $|0\rangle$ that is annihilated by all single-particle annihilation operators c_j – the state $|0\rangle$ is usually called the vacuum state of the Fock space. Formally, the fermionic Fock space of a given single-particle Hilbert space $\mathcal{H}$ is given by the exterior algebra

$$\mathcal{F}(\mathcal{H}) = \bigoplus_{n=1}^{\infty} \mathcal{H}^{\wedge n}, \quad (5.139)$$

where $\mathcal{H}^{\wedge n}$ denotes the antisymmetrised n -fold tensor power of $\mathcal{H}$. The algebra of operators on $\mathcal{F}(\mathcal{H})$ is generated by the fermionic creation and annihilation operators $c_j^\dagger$ and c_j satisfying

$$\{c_j, c_k^\dagger\} = \delta_{jk} \quad \text{and} \quad \{c_j, c_k\} = \{c_j^\dagger, c_k^\dagger\} = 0 \quad (5.140)$$

for all $j, k = 1, \dots, d$. One can combine the creation and annihilation operators into Hermitian operators

$$\gamma_{j1} = c_j^\dagger + c_j \quad \text{and} \quad \gamma_{j2} = i(c_j^\dagger - c_j), \quad (5.141)$$

that satisfy

$$\{\gamma_{js}, \gamma_{kr}\} = \{\gamma_{js}, \gamma_{kr}^\dagger\} = \{\gamma_{js}^\dagger, \gamma_{kr}^\dagger\} = 2\delta_{js, kr}. \quad (5.142)$$

These operators still generate the algebra of operators on $\mathcal{F}(\mathcal{H})$. Furthermore, they still provide a basis for Nambu space $\mathbf{H}$ due to the Nambu redundancy. We will come back to this shortly. The anticommutation relations shown in Eq. (5.142) correspond to the anticommutation relations characterising the real Clifford algebra $\text{Cl}_{2d}(\mathbb{R})$ of $2d$ generators, which illustrates that the operators on $\mathcal{F}(\mathcal{H})$ form an algebra that is isomorphic to $\text{Cl}_{2d}(\mathbb{R})$. A Bogoliubov vacuum can then be defined by

$$|B\rangle = \mathcal{R}(B) |\text{ref}\rangle = e^{\rho(X)} |\text{ref}\rangle, \quad (5.143)$$

where $\mathcal{R}(B)$ denotes a unitary representation of a given proper Bogoliubov matrix $B \in \text{SO}(2d) < \text{O}(2d)$ and where $|\text{ref}\rangle$ is some reference state. The second equality expresses B through the exponential map and uses the relation

$$\mathcal{R}(B) = e^{\rho(X)} \quad \text{for} \quad B = e^X \quad (5.144)$$

between a unitary representation $\mathcal{R}$ of $\text{SO}(2d)$ and the corresponding representation ρ of $\mathfrak{o}(2d)$ on the Fock space. The orthogonal group $\text{O}(2d)$ has a double-valued unitary representation

$$\begin{aligned} \mathcal{R} : \text{SO}(2d) &\rightarrow \text{Spin}(2d) \subset \text{Cl}_{2d}(\mathbb{R}) \\ G &\mapsto \pm \mathcal{R}(G) \end{aligned} \quad (5.145)$$

on the fermionic Fock space $\mathcal{F}(\mathcal{H})$ that is known as the spin-representation of $\text{O}(2d)$. As indicated in Eq. (5.145), the double-valuedness of the spin representation means that the spin representation maps every $G \in \text{O}(2d)$ to a pair of Fock space operators $\pm \mathcal{R}(G) \in \text{Spin}(2d)$. Neergård shows that the Fock space representation of the Bogoliubov group is precisely this spin representation [87]. He chooses the elementary fermion vacuum $|0\rangle$ as the reference state and defines the Bogoliubov vacuum as

$$|B\rangle = \pm \mathcal{R}(B) |0\rangle = \pm e^{\rho(X)} |0\rangle, \quad (5.146)$$

⁶Given that the exponential map is invertible at B .

where $\mathcal{R}(B)$ now denotes the spin representation of a proper Bogoliubov matrix $B \in SO(2d) < O(2d)$ and where $\rho(X)$ is the spin representation of the Lie algebra element $X \in \mathfrak{o}(2d)$ generating B via the exponential map. The sign ambiguity in Eq. (5.146) is a direct consequence of the double-valuedness of Eq. (5.145). For $B \in SO(2d)$, he arrives at the evocative form

$$|B\rangle = \pm \exp \left[\frac{1}{2} \Psi^\dagger \log(B) \Psi \right] |0\rangle, \quad (5.147)$$

where Ψ is the same Nambu spinor as before. When the reference state $|\text{ref}\rangle = |0\rangle$ is fixed, the double-valuedness of the spin representation translates to a double-valuedness of the Bogoliubov vacuum $|B\rangle$ that may cause sign problems in overlap formulas between BdG Fock states that require taking the square root. In Ref. [87], Neergård demonstrates in detail how this understanding of the Bogoliubov vacuum state can be used to determine and explain many of the established BdG overlap formulas.

Elegant as it is, the above formalism is not always the most convenient in practice. The main reason for this is that the exponential involves non-commuting field operators, which generally require careful tracking using a suitable normal ordering relation. Thus, practical applications often require simplified expressions. In the following, we will present the two most prominent variants known as the product state and the Thouless state. The product state can be defined as

$$|0\rangle_b^p = \prod_{j=1}^d b_j |0\rangle, \quad (5.148)$$

where $|0\rangle$ denotes the elementary fermion vacuum determined by $c_j |0\rangle \stackrel{!}{=} 0$. In this form, the state is not normalised and its norm $\mathcal{N}_p$ is given by

$$\mathcal{N}_p = (-1)^{\frac{d(d-1)}{2}} \text{Pf} \begin{pmatrix} V^\top U & VV^* \\ -V^\dagger V & U^\dagger V^* \end{pmatrix}, \quad (5.149)$$

which we will discuss in more detail shortly. The nice thing about the product state is that its form makes it obvious that it satisfies Eq. (5.137). However, this simplicity does not come for free. As we will soon see, the product state as defined in Eq. (5.148) can only exist when the matrix V of the underlying Bogoliubov transformation B is invertible. The other simplified variant of the Bogoliubov vacuum we mentioned is the Thouless state. It can be defined as

$$|0\rangle_b^T = \exp \left[\frac{1}{2} (\mathbf{c}^\dagger)^\top S \mathbf{c}^\dagger \right] |0\rangle, \quad (5.150)$$

where

$$S = (VU^{-1})^* \quad (5.151)$$

is an auxiliary skew-symmetric matrix. The definition of S immediately shows that the Thouless state comes with a similar limitation: it is only defined when U is invertible. The norm of $|0\rangle_b^T$ is

$$\mathcal{N}_T = (-1)^{\frac{d(d+1)}{2}} \text{Pf} \begin{pmatrix} S & -\mathbb{1}_d \\ \mathbb{1}_d & -S^* \end{pmatrix}, \quad (5.152)$$

which we will derive along with the norm of the product state in the next section. Compared to the product state, it is less straightforward to recognise the Thouless state as a Bogoliubov vacuum. Therefore, we provide a proof of its validity as a Bogoliubov vacuum state in App. A.3. Furthermore, a quick proof of the skew-symmetry of the auxiliary matrix S is presented in App. A.4.

Note that the form of the Thouless state $|0\rangle_b^T$ is reminiscent of the formal Bogoliubov vacuum $|B\rangle$ given in Eq. (5.147). Indeed, Ref. [87] states the relation

$$|0\rangle_b^T = \frac{|B\rangle}{\langle 0|B\rangle} \quad (5.153)$$

between them. Still, there are a few important differences between $|0\rangle_b^T$ and $|B\rangle$ that we have to be aware of. The vacuum state $|B\rangle$ is defined by a unitary exponential operator, which is constructed according

to quite general principles as the spin representation of $O(2d)$ on the BdG Fock space. As a result, it inherits a sign ambiguity from the double-valuedness of the spin representation, as shown in Eqs. (5.145) and (5.147). In contrast, the Thouless state arises from a non-unitary exponential operator, which only exists when U is invertible. However, if it *does* exist, the Thouless state has no sign ambiguity because the arbitrary sign of $|B\rangle$ cancels out in the division, as evident from Eq. (5.153). Another thing that is clear from Eq. (5.153) is that $|0\rangle_b^T$ is only defined when $\langle 0|B\rangle \neq 0$. Its uniqueness therefore comes at the expense of generality. One may wonder whether $\langle 0|B\rangle \neq 0$ and the invertibility of U are two independent conditions. This is not the case. As we shall discuss later on, the overlap between $|B\rangle$ and $|0\rangle$ is given by

$$\langle 0|B\rangle = \sqrt{\det(U)}, \quad (5.154)$$

showing that $\langle 0|B\rangle \neq 0$ and the invertibility of U capture the same underlying constraint. On a more practical note, the formal Bogoliubov vacuum $|B\rangle$ is based on an operator exponential with non-commuting operators, so that it generally requires a normal ordering procedure. In contrast, the operator exponential defining $|0\rangle_b^T$ is built from pairs of fermionic creation operators, which commute among themselves, eliminating the need for normal ordering. If both the product state *and* the Thouless state exist, they are related by

$$|0\rangle_b^P = \text{Pf}(U^\dagger V^*) |0\rangle_b^T. \quad (5.155)$$

A proof of this relation is included in App. A.6. However, recall that $|0\rangle_b^P$ and $|0\rangle_b^T$ can only exist simultaneously when both U and V are invertible. This is not always the case. In fact, there are quite general circumstances that prevent U and V from being invertible. To find these, we first exploit the unitarity

$$B^\dagger B = \begin{pmatrix} U^\dagger & V^\dagger \\ V^\dagger & U^\dagger \end{pmatrix} \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix} = \begin{pmatrix} \mathbb{1}_d & 0 \\ 0 & \mathbb{1}_d \end{pmatrix} = \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix} \begin{pmatrix} U^\dagger & V^\dagger \\ V^\dagger & U^\dagger \end{pmatrix} = BB^\dagger \quad (5.156)$$

of B to arrive at a collection

$$\begin{aligned} U^\dagger U + V^\dagger V &= \mathbb{1}_d = UU^\dagger + V^* V^\dagger \\ U^\dagger V^* + V^\dagger U^* &= 0 = UV^\dagger + V^* U^\dagger \\ V^\dagger U + U^\dagger V &= 0 = VU^\dagger + U^* V^\dagger \\ V^\dagger V^* + U^\dagger U^* &= \mathbb{1}_d = VV^\dagger + U^* U^\dagger \end{aligned} \quad (5.157)$$

of conditions for the matrices U and V . Using Schur's determinant identity

$$\det \begin{pmatrix} A & B \\ C & D \end{pmatrix} = \det(A) \det(D - CA^{-1}B), \quad (5.158)$$

we can then show that

$$\begin{aligned} \det(B) &= \det(U) \det(U^* - VU^{-1}V^*) \\ &\stackrel{(\diamond)}{=} \det(U) \det(U^* + VV^\dagger U^{-1}U^\dagger) \\ &\stackrel{(\star)}{=} \det(U) \det(U^* + (\mathbb{1}_d - U^* U^\dagger)U^{-1}U^\dagger) \\ &= \det(U) \det(U^{-1}U^\dagger) \\ &= 1, \end{aligned} \quad (5.159)$$

where we have used $U^{-1}V^* = - (U^{-1}V^*)^\dagger$ for $(\diamond)$ and $VV^\dagger = \mathbb{1}_d - U^* U^\dagger$ for $(\star)$. Both of these can be obtained from Eqs. (5.157). Concretely, $(\star)$ is a rearranged version of the bottom right equation, while $(\diamond)$ is obtained by multiplying $V^* U^\dagger + UV^\dagger = 0$ by $U^\dagger^{-1}$ from the right and by U^{-1} from the left, so that

$$U^{-1}V^* = -V^\dagger U^{-1}U^\dagger = - (U^{-1}V^*)^\dagger. \quad (5.160)$$

Equation (5.159) tells us that U invertible $\implies \det(B) = 1$. The contraposition of this statement is $\det(B) \neq 1 \implies U$ not invertible, which, given that the only possible determinants of Bogoliubov transformations are $\det(B) = \pm 1$, can be refined

$$\det(B) = -1 \text{ prevents } U \text{ from being invertible.}$$

	$\det(B) = +1$	$\det(B) = -1$
d even	—	$\det(U) = \det(V) = 0$
d odd	$\det(V) = 0$	$\det(U) = 0$

Table 5.1: Invertibility conditions for U and V . Entries indicate which matrices are forced to be singular.

To get a similar statement for V we note that Eqs. (5.157) are invariant under an exchange of U and V , whereas the determinant of B transforms as

$$\det(B) = \det \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix} = (-1)^d \det \begin{pmatrix} V & U^* \\ U & V^* \end{pmatrix} =: (-1)^d \det(B') \quad (5.161)$$

under the same exchange. For odd d , we can therefore map an improper Bogoliubov matrix B with $\det(B) = -1$ to a proper Bogoliubov matrix B' with $\det(B') = 1$ by exchanging U and V . Since Eqs. (5.157) are left invariant under this exchange, we can repeat the above argument for B' and find that V invertible $\implies \det(B') = 1$. If we plug in the odd- d relation $\det(B') = -\det(B)$ from Eq. (5.161) backwards, we obtain $\det(B) \neq 1 \implies V$ not invertible, showing that

$$\det(B) = -1 \text{ prevents } V \text{ from being invertible when } d \text{ is odd.}$$

When d is even, the exchange of U and V does not change the determinant sign and $\det(B) = -1$ obstructs invertibility of both U and V . A summary of these constraints is given in Tab. 5.1. These limitations are quite severe. For odd d , the product state and the Thouless state can never coexist, while for even d , it is possible that neither state exists. In addition to these categorical restrictions, U and V may always become singular by coincidence. So even though the ground state of a given BdG Hamiltonian always exists, the simple Bogoliubov vacuum form in Eq. (5.136) is not going to be available in every situation.

This is a problem. Fortunately, there is a quite elegant solution. To appreciate it, we must first understand the “physical” reason why the construction of the product and Thouless state may fail. That reason, as it turns out, is an accidental annihilation of the reference state. Generally, the construction of every quasiparticle vacuum $|0\rangle_b$ as defined via Eq. (5.137) ultimately amounts to a systematic removal of b_j quasiparticles from some reference reference state $|\text{ref}\rangle$. The product state and the Thouless state use the elementary vacuum $|0\rangle$ of the c -fermions as a reference state. This is most apparent in the definition of the product state in Eq. (5.148), where we take the elementary fermion vacuum $|0\rangle$ and remove all Bogoliubov quasiparticles from it by successively acting on it with every Bogoliubov annihilation operator. The key observation is that

$$\prod_{j=1}^d b_j |0\rangle = 0 \quad (5.162)$$

means that there is at least one Bogoliubov quasiparticle mode, which is already empty in $|0\rangle$. It seems natural to simply exclude these empty modes from the product – after all, they are *already* empty – and proceed as usual. The problem with this strategy is, that it is not at all obvious, *which* Bogoliubov modes are empty in $|0\rangle$. We would require a systematic way of identifying empty Bogoliubov modes to make this work. In a favourable turn of events, there exists a decomposition algorithm that accomplishes exactly this. The Bloch–Messiah decomposition (BMD) may be described as a singular value decomposition (SVD) that is compatible with the particle-hole conjugation structure of Bogoliubov transformations [88]. For a given $2d \times 2d$ Bogoliubov transformation matrix B , it is a factorisation

$$B = \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix} = \begin{pmatrix} C^\dagger & 0 \\ 0 & C^\intercal \end{pmatrix} \begin{pmatrix} \bar{U} & \bar{V} \\ \bar{V} & \bar{U} \end{pmatrix} \begin{pmatrix} D & 0 \\ 0 & D^* \end{pmatrix} = \mathcal{C} \bar{B} \mathcal{D}^\dagger \quad (5.163)$$

of B into a product of a block diagonal unitary $2d \times 2d$ matrix $\mathcal{C}$, which is defined in terms of a unitary $d \times d$ matrix C , a real $2d \times 2d$ matrix $\bar{B}$, which is composed of two real $d \times d$ matrices $\bar{U}$ and $\bar{V}$, and another unitary $2d \times 2d$ matrix $\mathcal{D}^\dagger$, which is determined by a unitary $d \times d$ matrix D . The two $d \times d$

blocks of $\bar{B}$ are given by the diagonal and block-diagonal matrices

$$\bar{U} = \begin{pmatrix} \mathbb{0}_F & & \\ & \bigoplus_{p=1}^{d_P} u_p \mathbb{1}_2 & \\ & & \mathbb{1}_E \end{pmatrix} \quad \text{and} \quad \bar{V} = \begin{pmatrix} \mathbb{1}_F & & \\ & \bigoplus_{p=1}^{d_P} i v_p \sigma_y & \\ & & \mathbb{0}_E \end{pmatrix}. \quad (5.164)$$

Here, we have introduced three index sets

$$I_F := \{1, \dots, d_F\}, \quad I_P := \{1, \bar{1}, \dots, d_P, \bar{d}_P\}, \quad I_E := \{1, \dots, d_E\}, \quad (5.165)$$

that partition the original index $J = \{1, \dots, d\}$ set as

$$\begin{aligned} I_F &= \{1, \dots, d_F\} \mapsto \{1, \dots, d_F\} =: J_F \subset J \\ I_P &= \{1, \bar{1}, \dots, d_P, \bar{d}_P\} \mapsto \{d_F + 1, d_F + 2, \dots, d_F + 2d_P - 1, d_F + 2d_P\} =: J_P \subset J \\ I_E &= \{1, \dots, d_E\} \mapsto \{d_F + 2d_P + 1, \dots, d\} =: J_E \subset J, \end{aligned} \quad (5.166)$$

and we write

$$J = J_F \cup J_P \cup J_E. \quad (5.167)$$

Note that we have $d = d_F + 2d_P + d_E$ by construction. The new index sets are named in anticipation of the so-called filled, paired, and empty states to which they correspond. This will become clear shortly. The cardinalities

$$|I_F| = |J_F| = d_F, \quad |I_P| = |J_P| = 2d_P, \quad |I_E| = |J_E| = d_E \quad (5.168)$$

are given in terms of the “dimensions” d_F , d_P , and d_E of the filled, paired, and empty BMD sectors, respectively. Note that we will usually sum over $p = 1, \dots, d_P$ and implicitly account for the paired indices $\bar{p}$. For instance, in Eq. (5.164), the P paired blocks are 2×2 matrices in the paired indices. Each of these blocks is characterised by two numbers u_p and v_p that fulfil

$$u_p^2 + v_p^2 = 1 \quad \text{and} \quad u_p, v_p > 0 \quad (5.169)$$

for all $p = 1, \dots, d_P$. Finally, we wrote $\mathbb{0}_F$ and $\mathbb{0}_E$ for the d_F - and the d_E -dimensional zero matrix, and $\mathbb{1}_F$ and $\mathbb{1}_E$ for the d_F - and the d_E -dimensional unit matrix in Eq. (5.164). The actual construction of the BMD is rather lengthy, occasionally tedious, and often conspicuously absent in the literature. The interested reader may refer to App. A.5 for a comprehensive discussion.

It is worth mentioning that the form of the BMD in Eq. (5.163) differs from the standard form found in much of the literature, cf. e.g. Ref. [83]. The main difference is that Eq. (5.163) is asymmetric in the unitary matrices $\mathcal{C}$ and $\mathcal{D}$, combining $\mathcal{C}$ and $\mathcal{D}^\dagger$ instead of $\mathcal{C}$ and $\mathcal{D}$ or $\mathcal{C}^\dagger$ and $\mathcal{D}^\dagger$. The reason why we define $\mathcal{C}$ and $\mathcal{D}^\dagger$ in this way is because it makes the *following* definitions more symmetrical. Concretely, if we plug Eq. (5.163) into the definition of the Bogoliubov–Nambu spinor from Eq. (5.35), we get

$$\begin{pmatrix} D & 0 \\ 0 & D^* \end{pmatrix} \begin{pmatrix} \mathbf{b} \\ \mathbf{b}^\dagger \end{pmatrix} = \begin{pmatrix} \bar{U}^\tau & \bar{V}^\tau \\ \bar{V}^\tau & \bar{U}^\tau \end{pmatrix} \begin{pmatrix} C & 0 \\ 0 & C^* \end{pmatrix} \begin{pmatrix} \mathbf{c} \\ \mathbf{c}^\dagger \end{pmatrix}, \quad (5.170)$$

which, using $\bar{U}^\tau = \bar{U}$, takes the simple form

$$\begin{pmatrix} \bar{\mathbf{b}} \\ \bar{\mathbf{b}}^\dagger \end{pmatrix} = \begin{pmatrix} \bar{U} & \bar{V}^\tau \\ \bar{V}^\tau & \bar{U} \end{pmatrix} \begin{pmatrix} \bar{\mathbf{c}} \\ \bar{\mathbf{c}}^\dagger \end{pmatrix} \quad (5.171)$$

upon defining new fermionic (quasi-)particle modes

$$\bar{c}_j = \sum_{m=1}^d C_{jm} c_m \quad \text{and} \quad \bar{b}_j = \sum_{m=1}^d D_{jm} b_m. \quad (5.172)$$

Note that this transformation only mixes the elementary and Bogoliubov annihilation operators amongst themselves, so the new fermion modes have the same vacuum as the original ones, i.e.

$$\bar{c}_j |0\rangle = 0 \quad \text{and} \quad \bar{b}_j |0\rangle_b = 0 \quad (5.173)$$

for all $j = 1, \dots, d$. The new elementary fermion annihilators $\bar{c}_n$ correspond to a basis of single-particle states that is known as the “canonical” single-particle basis because it diagonalises the one-particle reduced density matrix ρ [83]. While this is quite useful in its own right, our primary interest lies in the new Bogoliubov quasiparticle operators. Specifically, the (block-)diagonal structure of the matrices $\bar{U}$ and $\bar{V}$ from Eq. (5.164) allows us to identify three types of new Bogoliubov quasiparticles:

- i) the $f = 1, \dots, d_F$ “filled” modes, where $v_f = 1$ and $u_f = 0$, such that

$$\bar{b}_f = \bar{c}_f^\dagger. \quad (5.174)$$

- ii) the $p = 1, \dots, d_P$ “paired” modes, where $u_p, v_p > 0$, such that

$$\bar{b}_p = u_p \bar{c}_p - v_p \bar{c}_p^\dagger \quad \text{and} \quad \bar{b}_{\bar{p}} = u_p \bar{c}_{\bar{p}} + v_p \bar{c}_{\bar{p}}^\dagger. \quad (5.175)$$

- iii) the $e = 1, \dots, d_E$ “empty” modes, where $v_e = 0$ and $u_e = 1$, such that

$$\bar{b}_e = \bar{c}_e. \quad (5.176)$$

If we plug the inverse transformation $b_j = \sum_{n=1}^d D_{jn}^\dagger \bar{b}_n$ of the $\bar{b}_n$ from Eq. (5.172) into Eq. (5.148) we end up with

$$\begin{aligned} |0\rangle_b^P &= \prod_{j=1}^d b_j |0\rangle \\ &= \prod_{j=1}^d \left(\sum_{n_j} D_{jn_j}^\dagger \bar{b}_{n_j} \right) |0\rangle \\ &= \sum_{n_1=1}^d \cdots \sum_{n_d=1}^d \left(\prod_{j=1}^d D_{jn_j}^\dagger \bar{b}_{n_j} \right) |0\rangle \\ &\stackrel{(\diamond)}{=} \left[\sum_{\pi \in S_d} \text{sign}(\pi) \left(\prod_{j=1}^d D_{j\pi(j)}^\dagger \right) \right] \bar{b}_1 \cdots \bar{b}_d |0\rangle \\ &\stackrel{(*)}{=} \det(D^\dagger) \prod_{j=1}^d \bar{b}_j |0\rangle. \end{aligned} \quad (5.177)$$

In $(\diamond)$ we exploited that $\bar{b}_j^2 = 0$ for all j so that

$$\sum_{n_1=1}^d \cdots \sum_{n_d=1}^d \left(\prod_{j=1}^d D_{jn_j}^\dagger \bar{b}_{n_j} \right) \quad (5.178)$$

only contains terms where each $\bar{b}_j$ appears exactly once. As a result, the whole expression reduces to a sum over all permutations of operator products, i.e.

$$\sum_{n_1=1}^d \cdots \sum_{n_d=1}^d \left(\prod_{j=1}^d D_{jn_j}^\dagger \bar{b}_{n_j} \right) \longrightarrow \sum_{\pi \in S_d} D_{1\pi(1)}^\dagger \cdots D_{d\pi(d)}^\dagger \bar{b}_{\pi(1)} \cdots \bar{b}_{\pi(d)}, \quad (5.179)$$

where $\pi(i)$ denotes the image of the index i under the permutation $\pi \in S_d$ of the symmetric group of d elements. Using $\{\bar{b}_j, \bar{b}_k\} = 0$ we can bring each permutation in a fixed reference order if we account for the anticommutator sign $\bar{b}_j \bar{b}_k = -\bar{b}_k \bar{b}_j$ that comes with every transposition. We chose the natural reference order $\bar{b}_1 \cdots \bar{b}_d$, with which we get

$$\left[\sum_{\pi \in S_d} \text{sign}(\pi) D_{1\pi(1)}^\dagger \cdots D_{d\pi(d)}^\dagger \right] \bar{b}_1 \cdots \bar{b}_d, \quad (5.180)$$

where we took the operator string out of the sum because it no longer depends on the summation over permutations. The convenient choice of reference order means that the anticommutator signs are precisely the permutation signs $\text{sign}(\pi) = (-1)^{N_\pi}$ with the number N_π of neighbouring transpositions generating the elements $\pi \in S_d$ of the d -th permutation group. This allows us to plug in the definition

$$\det(D^\dagger) := \sum_{\pi \in S_d} \text{sign}(\pi) D_{1\pi(1)}^\dagger \cdots D_{d\pi(d)}^\dagger \quad (5.181)$$

of the determinant of $D^\dagger$ in $(\star)$ and get the final expression. Now, D is a unitary matrix so $\det(D^\dagger)$ is only a phase factor. This shows that $|0\rangle_b$ and $|0\rangle_{\bar{b}}$ do indeed correspond to the same quantum state. For now, we disregard this additional phase factor and use the definitions of the empty and filled Bogoliubov quasiparticle modes to write

$$\begin{aligned} |0\rangle_b^p &\propto \prod_{j=1}^d \bar{b}_j |0\rangle \\ &= \prod_{f=1}^{d_F} \bar{b}_f \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_{\bar{p}} \prod_{e=1}^{d_E} \bar{b}_e |0\rangle \\ &= \prod_{f=1}^{d_F} \bar{c}_f^\dagger \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_{\bar{p}} \prod_{e=1}^{d_E} \bar{c}_e |0\rangle. \end{aligned} \quad (5.182)$$

Clearly, this expression vanishes whenever $d_E > 0$. However, we already knew this, because Eq. (5.164) tells us that $d_E > 0$ makes $\bar{V}$ and hence V singular, giving the same condition as before. What we did *not* know before was how to isolate and remove the troublesome empty Bogoliubov modes from the equation. This is what Eq. (5.182) allows us to do. We can simply leave out the empty Bogoliubov modes and define the truncated product state as

$$|\bar{0}\rangle_b^p := \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_{\bar{p}} \prod_{f=1}^{d_F} \bar{c}_f^\dagger |0\rangle, \quad (5.183)$$

where we moved the filled modes to the right, using that $\bar{c}_f^\dagger = \bar{b}_f$ and $\{\bar{b}_j, \bar{b}_k\} = 0$ so that commuting the $\bar{c}_f^\dagger$ past the even number of paired Bogoliubov annihilation operators $\prod_p \bar{b}_p \bar{b}_{\bar{p}}$ does not yield an extra sign. The truncation of the product state can be reinterpreted in terms of the reference state. Namely, the truncation procedure is equivalent to choosing the reference state $|\text{ref}\rangle$ as

$$|\text{ref}\rangle := |0'\rangle = \prod_{e=1}^{d_E} \bar{b}_e^\dagger |0\rangle, \quad (5.184)$$

in which the Bogoliubov modes that are empty in $|0\rangle$ are manually occupied beforehand. With this, the product state becomes

$$\begin{aligned} |0'\rangle_b^p &= \prod_{f=1}^{d_F} \bar{b}_f \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_{\bar{p}} \prod_{e=1}^{d_E} \bar{b}_e |0'\rangle \\ &= \prod_{f=1}^{d_F} \bar{b}_f \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_{\bar{p}} \prod_{e=1}^{d_E} \bar{b}_e \prod_{e=1}^{d_E} \bar{b}_e^\dagger |0\rangle \\ &= \varepsilon_E \prod_{f=1}^{d_F} \bar{b}_f \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_{\bar{p}} \prod_{e=1}^{d_E} \bar{b}_e \bar{b}_e^\dagger |0\rangle \\ &= \varepsilon_E \prod_{f=1}^{d_F} \bar{b}_f \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_{\bar{p}} \prod_{e=1}^{d_E} \left(1 - \cancel{\bar{b}_e^\dagger \bar{b}_e}\right) |0\rangle \\ &= \varepsilon_E \prod_{f=1}^{d_F} \bar{b}_f \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_{\bar{p}} |0\rangle \\ &= \varepsilon_E |\bar{0}\rangle_b^p, \end{aligned} \quad (5.185)$$

where we defined the combinatorial sign

$$\varepsilon_E = (-1)^{\frac{d_E(d_E-1)}{2}}, \quad (5.186)$$

that results from the $d_E(d_E - 1)/2$ transpositions that implement the permutation

$$\bar{b}_1 \cdots \bar{b}_{d_E} \bar{b}_1^\dagger \cdots \bar{b}_{d_E}^\dagger \longrightarrow \bar{b}_1 \bar{b}_1^\dagger \cdots \bar{b}_{d_E} \bar{b}_{d_E}^\dagger, \quad (5.187)$$

and cancelled all terms with $\bar{b}_e$ -annihilation operators on the right because $\bar{b}_e |0\rangle = 0$ for the empty Bogoliubov modes. Up to an irrelevant global sign, this reproduces the truncated product state.

The key advantage of Eq. (5.183) is that it is well-defined and constitutes a ground state of the underlying BdG Hamiltonian, regardless of whether V is singular ($d_E > 0$) or not ($d_E = 0$). Finally, we note that one can bring $|\bar{0}\rangle_b^P$ into the practical form

$$|\bar{0}\rangle_b^P =: \mathcal{N}_P \prod_{p=1}^{d_P} \left(u_p + v_p \bar{c}_p^\dagger \bar{c}_p^\dagger \right) \prod_{f=1}^{d_F} \bar{c}_f^\dagger |0\rangle \quad (5.188)$$

by plugging in the paired Bogoliubov modes from Eq. (5.175) and rewriting their product as

$$\begin{aligned} \prod_{p=1}^{d_P} \bar{b}_p \bar{b}_p &= \prod_{p=1}^{d_P} \left(u_p \bar{c}_p - v_p \bar{c}_p^\dagger \right) \left(u_p \bar{c}_p + v_p \bar{c}_p^\dagger \right) \\ &= \prod_{p=1}^{d_P} \left(\cancel{u_p^2 \bar{c}_p \bar{c}_p} + u_p v_p \bar{c}_p \bar{c}_p^\dagger - v_p u_p \cancel{\bar{c}_p^\dagger \bar{c}_p^\dagger} - v_p^2 \bar{c}_p^\dagger \bar{c}_p^\dagger \right) \\ &= \prod_{p=1}^{d_P} \left(u_p v_p \left(1 - \cancel{\bar{c}_p \bar{c}_p} \right) + v_p^2 \bar{c}_p^\dagger \bar{c}_p^\dagger \right) \\ &= \prod_{p=1}^{d_P} v_p \left(u_p + v_p \bar{c}_p^\dagger \bar{c}_p^\dagger \right) \\ &=: \mathcal{N}_P \prod_{p=1}^{d_P} \left(u_p + v_p \bar{c}_p^\dagger \bar{c}_p^\dagger \right), \end{aligned} \quad (5.189)$$

where we defined the truncated norm factor

$$\mathcal{N}_P := \prod_{p=1}^{d_P} v_p \quad (5.190)$$

and cancelled all the terms with $\bar{c}$ -annihilation operators on the right because the whole expression acts on the $\bar{c}$ -vacuum $|0\rangle$ afterwards. Thus, we obtain a normalised truncated product state

$$|\bar{0}\rangle_b^P = \prod_{p=1}^{d_P} \left(u_p + v_p \bar{c}_p^\dagger \bar{c}_p^\dagger \right) \prod_{f=1}^{d_F} \bar{c}_f^\dagger |0\rangle, \quad (5.191)$$

which is distinctly reminiscent of the BCS vacuum state we stated in Eq. (5.13), although the paired indices do not (necessarily) refer to quasi-momentum. Based on the normalised truncated product form of the quasiparticle vacuum from Eq. (5.191) we can construct the excited energy eigenstates as BdG Fock states as

$$|n_1, \dots, n_d\rangle_b := \prod_{m=1}^d (b_m^\dagger)^{n_m} |\bar{0}\rangle_b^P, \quad (5.192)$$

where $n_m = 0, 1$ denotes the occupation of the m -th quasiparticle state.

5.4 Bogoliubov Overlaps

For many practical purposes, we are interested in overlaps between BdG Fock states. It turns out that finding a reliable overlap formula is a surprisingly hard problem. In Ref. [87], Neergård demonstrated that this is related to the fact that the Fock space representation of Bogoliubov transformations corresponds to the spin representation of an orthogonal group. The inherent double-valuedness of the spin representations then leads to a sign-ambiguity for all overlap formulas involving a square root. Let us go through this in more detail.

The simplest BdG overlaps of interest are overlaps between different vacua. This includes things like the overlaps $\langle 0|0\rangle_b^P$ and $\langle 0|0\rangle_b^T$ between the Bogoliubov vacua $|0\rangle_b^P$ and $|0\rangle_b^T$ and their reference vacuum $|0\rangle$, but also norm squares like ${}_b^P\langle 0|0\rangle_b^P$ and ${}_b^T\langle 0|0\rangle_b^T$. Most BdG overlap formulas were first developed for the Thouless state. For example, Ref. [83] presents the Onishi formula

$$\langle 0|\tilde{0}\rangle_b^T = \sqrt{\det(U^*)} \quad (5.193)$$

for the overlap between a renormalised Thouless state

$$|\tilde{0}\rangle_b^T = \langle 0|\tilde{0}\rangle_b^T \exp\left[\frac{1}{2}(\mathbf{c}^\dagger)^\top S \mathbf{c}^\dagger\right] |0\rangle \quad (5.194)$$

and the elementary vacuum, and the expression

$${}_b^T\langle \tilde{0}'|\tilde{0}\rangle_b^T = \sqrt{\det(U^\dagger U' + V^\dagger V')} \quad (5.195)$$

for the overlap between two distinct renormalised Thouless states

$$|\tilde{0}\rangle_b^T = \langle 0|\tilde{0}\rangle_b^T \exp\left[\frac{1}{2}(\mathbf{c}^\dagger)^\top S \mathbf{c}^\dagger\right] |0\rangle \quad \text{and} \quad |\tilde{0}'\rangle_b^T = \langle 0|\tilde{0}'\rangle_b^T \exp\left[\frac{1}{2}(\mathbf{c}^\dagger)^\top S' \mathbf{c}^\dagger\right] |0\rangle. \quad (5.196)$$

Above, the matrices S and S' are given by $S = (VU^{-1})^*$ and $S' = (V'U'^{-1})^*$, respectively. Equation (5.195) is sometimes called an Onishi formula in the literature as well. Both of the Onishi formulas above refer to a paper [89] by Onishi and Yoshida, in which they use a Thouless state of the form

$$|0\rangle_b^T = \frac{1}{\langle 0|\tilde{0}\rangle_b^T} |\tilde{0}\rangle_b^T = \exp\left[\frac{1}{2}(\mathbf{c}^\dagger)^\top S \mathbf{c}^\dagger\right] |0\rangle \quad (5.197)$$

with $\langle 0|0\rangle_b^T = 1$, and presented the overlap formula

$$\begin{aligned} {}_b^T\langle 0'|0\rangle_b^T &= \exp\left[\frac{1}{2}\text{tr}(\log(1 + S'^\dagger S))\right] \\ &= \sqrt{\det(1 + S'^\dagger S)}. \end{aligned} \quad (5.198)$$

Note that Eq. (5.198) can be used to reproduce Eq. (5.195), since

$$\begin{aligned} {}_b^T\langle \tilde{0}'|\tilde{0}\rangle_b^T &\stackrel{(\diamond)}{=} {}_b^T\langle \tilde{0}'|0\rangle {}_b^T\langle 0|\tilde{0}\rangle_b^T \\ &= \sqrt{\det(U')} \sqrt{\det(U^*)} \sqrt{\det(1 + S'^\dagger S)} \\ &\stackrel{(\star)}{=} \sqrt{\det(U')} \sqrt{\det(U^*)} \sqrt{\det(1 + U'^\top \tau^{-1} V'^\top \tau V^* U^{*-1})} \\ &= \sqrt{\det(U')} \sqrt{\det(U^*)} \sqrt{\det(U'^\top \tau^{-1})} \sqrt{\det(U'^\top \tau U^* + V'^\top \tau V^*)} \sqrt{\det(U^{*-1})} \\ &\stackrel{(*)}{=} \frac{\sqrt{\det(U')} \sqrt{\det(U^*)}}{\sqrt{\det(U'^\top \tau)} \sqrt{\det(U^*)}} \sqrt{\det(U'^\top \tau U^* + V'^\top \tau V^*)} \\ &\stackrel{(\Delta)}{=} \sqrt{\det(U^\dagger U' + V^\dagger V')}, \end{aligned} \quad (5.199)$$

where we plugged in Eq. (5.197) in $(\diamond)$, inserted the definitions of $S'^\dagger$ and S in $(\star)$, and finally used $\det(A^{-1}) = 1/\det(A)$ in $(*)$ as well as $\det(A^\top) = \det(A)$ in (Δ) .

The renormalised Thouless state in Eqs. (5.194) and (5.196) is essentially the formal Bogoliubov vacuum $|B\rangle$ from Eq. (5.147), but with the full Nambu-spinor exponential reduced to one involving only creation operators. The contractions of the mixed creation-annihilation terms required for this reduction precisely yield the normalisation factor in Eqs. (5.194) and (5.196). As a result, Eqs. (5.194) and (5.196) inherit a sign ambiguity from the double-valuedness of the spin representation of $O(2d)$. Neergård argues that this sign ambiguity manifests as the intrinsic sign ambiguity of the square root in Eqs. (5.193) and (5.195). In contrast, the Thouless state in Eq. (5.197) is unique because the arbitrary sign of $|0\rangle_b^T$ cancels out in the division by $\langle 0|0\rangle_b^T$. Of course, this is only possible when $\langle 0|0\rangle_b^T \neq 0$, in which case it is equivalent to fixing the sign of $|0\rangle_b^T$ by demanding $\langle 0|0\rangle_b^T \stackrel{!}{=} 1$. The sign ambiguity of the square root in Eq. (5.198) can then be resolved defining

$$|0(\tau)\rangle_b^T := \exp\left[\frac{\tau}{2} (\mathbf{c}^\dagger)^\top S \mathbf{c}^\dagger\right] |0\rangle \quad (5.200)$$

and requiring that

$${}_b^T \langle 0' | 0(\tau) \rangle_b^T = \sqrt{\det(1 + \tau S'^\dagger S)} \quad (5.201)$$

be a continuous function of $\tau \in [0, 1]$. Since the limit of $\tau \rightarrow 0$ is well-defined it can be used to choose a sign of the square root. This is not an option for the overlap formulas based on $|0\rangle_b^T \propto |B\rangle$ mentioned before. A comprehensive discussion of these ideas is given in Ref. [87].

Here, we are going to focus on an overlap formula that is due to Robledo [90]. In contrast to the above formulas, the Robledo formula does not require taking a square root. Instead, it is given by

$${}_b^T \langle 0' | 0 \rangle_b^T = (-1)^{\frac{d(d+1)}{2}} \text{Pf} \begin{pmatrix} S & -\mathbb{1}_d \\ \mathbb{1}_d & -S'^* \end{pmatrix}. \quad (5.202)$$

One way to derive this result is through fermionic coherent states, which allow us to exploit the underlying Grassmann algebra and the associated Berezin integral techniques. Specifically, we can write

$$\begin{aligned} {}_b^T \langle 0' | 0 \rangle_b^T &= \langle 0 | \exp\left[-\frac{1}{2} S'_{jk}^* c_j c_k\right] \exp\left[\frac{1}{2} S_{jk} c_j^\dagger c_k^\dagger\right] |0\rangle \\ &\stackrel{(\diamond)}{=} \langle 0 | \exp\left[-\frac{1}{2} S'_{jk}^* c_j c_k\right] \left(\int \left(\prod_{i=1}^d d\xi_i^* d\xi_i \right) \exp[-\xi_i^* \xi_i] |\xi\rangle \langle \xi| \right) \exp\left[\frac{1}{2} S_{jk} c_j^\dagger c_k^\dagger\right] |0\rangle \\ &= \int \left(\prod_{i=1}^d d\xi_i^* d\xi_i \right) \exp[-\xi_i^* \xi_i] \langle 0 | \exp\left[-\frac{1}{2} S'_{jk}^* c_j c_k\right] |\xi\rangle \langle \xi| \exp\left[\frac{1}{2} S_{jk} c_j^\dagger c_k^\dagger\right] |0\rangle \\ &\stackrel{(*)}{=} \int \left(\prod_{i=1}^d d\xi_i^* d\xi_i \right) \exp[-\xi_i^* \xi_i] \langle 0 | \xi \rangle \exp\left[-\frac{1}{2} S'_{jk}^* \xi_j \xi_k\right] \exp\left[\frac{1}{2} S_{jk} \xi_j^* \xi_k^*\right] \langle \xi | 0 \rangle \\ &\stackrel{(*)}{=} \int \left(\prod_{i=1}^d d\xi_i^* d\xi_i \right) \exp[-\xi_i^* \xi_i] \exp\left[-\frac{1}{2} S'_{jk}^* \xi_j \xi_k\right] \exp\left[\frac{1}{2} S_{jk} \xi_j^* \xi_k^*\right] \\ &\stackrel{(\triangle)}{=} (-1)^{\frac{d(d+1)}{2}} \int d\boldsymbol{\mu} \exp\left[\frac{1}{2} \boldsymbol{\mu}^\top M \boldsymbol{\mu}\right] \\ &\stackrel{(\bullet)}{=} (-1)^{\frac{d(d+1)}{2}} \text{Pf}(M), \end{aligned} \quad (5.203)$$

where we used an Einstein notation to improve readability: indices that appear twice are implicitly summed over unless they are explicitly included in a product. In $(\diamond)$, we inserted the completeness relation

$$\mathbb{1} = \int \left(\prod_{i=1}^d d\xi_i^* d\xi_i \right) \exp[-\xi_i^* \xi_i] |\xi\rangle \langle \xi| \quad (5.204)$$

of fermionic coherent states

$$|\xi\rangle := e^{-\xi_i c_i^\dagger} |0\rangle. \quad (5.205)$$

Note that $\langle 0|\xi\rangle = 1$ according to this definition. Furthermore, ξ_i and ξ_i^* are mutually conjugate complex Grassmann numbers, i.e. elements of the exterior algebra over the complex numbers satisfying

$$\{\xi_j, \xi_k\} = \{\xi_j^*, \xi_k^*\} = \{\xi_j^*, \xi_k\} = 0, \quad (5.206)$$

and, as a direct consequence, $\xi_j^2 = \xi_j^{*2} = 0$ for all $j, k = 1, \dots, d$. Next, we plugged in the eigenvalue equations

$$c_i |\xi\rangle = \xi_i |\xi\rangle \quad \text{and} \quad \langle \xi| c_i^\dagger = \xi_i^* \langle \xi| \quad (5.207)$$

of fermionic coherent states in $(\star)$ and used that $\langle 0|\xi\rangle = 1$ in $(*)$. Finally, we took advantage of the fact that the even products of the Grassmann numbers commute, which allowed us to combine the exponentials in $(*)$ and rewrite the resulting exponent as a vector-matrix-vector product between vectors

$$\boldsymbol{\mu} := (\xi_1^* \dots \xi_d^* \xi_1 \dots \xi_d)^\top \quad (5.208)$$

and a skew-symmetric matrix

$$M := \begin{pmatrix} S & -\mathbb{1}_d \\ \mathbb{1}_d & -S'^* \end{pmatrix} \quad (5.209)$$

in (Δ) . Additionally, we transformed the integral measure as

$$\prod_{i=1}^d d\xi_i^* d\xi_i = (-1)^{\frac{d(d+1)}{2}} d\xi_d \dots d\xi_1 d\xi_d^* \dots d\xi_1^* = (-1)^{\frac{d(d+1)}{2}} d\mu_{2d} \dots d\mu_1 =: (-1)^{\frac{d(d+1)}{2}} d\boldsymbol{\mu}, \quad (5.210)$$

where the sign prefactor is a consequence of the $d(d+1)/2$ transpositions that are necessary to rearrange the initial integral measure into the new one, i.e.

$$\begin{aligned} \prod_{i=1}^d d\xi_i^* d\xi_i &= (-1)^d \prod_{i=1}^d d\xi_i d\xi_i^* \\ &= (-1)^d d\xi_d d\xi_d^* \dots (d\xi_2 d\xi_2^*) d\xi_1 d\xi_1^* \\ &= (-1)^d d\xi_d d\xi_d^* \dots d\xi_1 (d\xi_2 d\xi_2^*) d\xi_1^* \\ &= (-1)^d d\xi_1 \dots d\xi_d d\xi_d^* \dots d\xi_1^* \\ &= (-1)^d (-1)^{\frac{d(d-1)}{2}} d\xi_d \dots d\xi_1 d\xi_d^* \dots d\xi_1^* \\ &= (-1)^{\frac{d(d+1)}{2}} d\xi_d \dots d\xi_1 d\xi_d^* \dots d\xi_1^*. \end{aligned} \quad (5.211)$$

Here, we first flipped all pairs of mutually conjugate differentials, which requires d transpositions and produces a sign factor of

$$N(\text{flip } d \text{ pairs}) = (-1)^d. \quad (5.212)$$

Then we successively moved the pairs of mutually conjugate Grassmann differentials to the proper position on the right as indicated between the second and third line of equation. Once more, this does not produce any signs because pairs of Grassmann objects commute with everything else. Finally, we reverse the order of $d\xi_1 \dots d\xi_d$ which requires $\sum_{k=1}^{d-1} k = d(d-1)/2$ transpositions and therefore yields an additional sign factor of

$$N(\text{reverse order}) = (-1)^{\frac{d(d-1)}{2}}, \quad (5.213)$$

giving the overall sign factor of

$$N(\text{flip } d \text{ pairs}) N(\text{reverse order}) = (-1)^d (-1)^{\frac{d(d-1)}{2}} = (-1)^{\frac{d(d+1)}{2}}. \quad (5.214)$$

in the final line of Eq. (5.211). We chose to explicitly rearrange the initial integral measure because it allows us to immediately plug in the following Berezin–Grassmann integration results in the final line $(\bullet)$

of Eq. (5.203): let $\theta_1, \dots, \theta_{2n}$ be the $2n$ generators of a Grassmann algebra and let A be a skew-symmetric $2n \times 2n$ matrix, then

$$\int d\theta_{2n} \dots d\theta_1 e^{\frac{1}{2} \theta_i A_{ij} \theta_j} = \text{Pf}(A). \quad (5.215)$$

The formula given in Eq. (5.202) determines the overlap between two Thouless Bogoliubov vacuum states without having to take a square root. It is natural to wonder whether we can find a similar result for the product form of the Bogoliubov vacuum. It turns out that this is possible using the relation Eq. (5.155) between the Thouless and the product state, with which Eq. (5.202) becomes

$${}_b \langle 0' | 0 \rangle_b^{\text{p}} = (-1)^{\frac{d(d+1)}{2}} \text{Pf} \begin{pmatrix} V'^{\dagger} U' & V'^{\dagger} V^* \\ -V^{\dagger} V' & U^{\dagger} V^* \end{pmatrix}. \quad (5.216)$$

The full calculation required to show this is given in App. A.7. The above Pfaffian formula for the overlap between product state representations of different Bogoliubov vacua can be generalised to overlaps between all BdG Fock states. This was done by Bertsch and Robledo [90–93], who showed that

$${}_b \langle n'_r | n_s \rangle_b = \pm \frac{1}{\mathcal{N} \mathcal{N}'} \text{Pf}(\mathbb{M}), \quad (5.217)$$

where

$$|n_s\rangle_b := b_{m_1}^{\dagger} \dots b_{m_s}^{\dagger} |\bar{0}\rangle_b^{\text{p}} \quad (5.218)$$

denotes a BdG Fock state with s excitations specified by the index set $s \equiv \{m_1, \dots, m_s\} \subset [1, N]$, and where

$$\mathbb{M} = \begin{pmatrix} \bar{V}'^{\dagger} \bar{U}' & \bar{V}'^{\dagger} C'^{\dagger} V_r'^* & \bar{V}'^{\dagger} C'^{\dagger} U_s & \bar{V}'^{\dagger} C'^{\dagger} C \bar{V} \\ \cdot & U_r'^{\dagger} V_r'^* & U_r'^{\dagger} U_s & U_r'^{\dagger} C \bar{V} \\ \cdot & \cdot & V_s^{\dagger} U_s & V_s^{\dagger} C \bar{V} \\ \cdot & \cdot & \cdot & \bar{U} \bar{V} \end{pmatrix}, \quad (5.219)$$

is a skew-symmetric matrix given in terms of the Bogoliubov matrices (U, V) associated to $|n_s\rangle_b$ and (U', V') associated to $|n'_r\rangle_b$ [93, 94]. The matrices $\bar{U}$, $\bar{V}$, C and $\bar{U}'$, $\bar{V}'$, C' denote the BMD components of (U, V) and (U', V') , respectively. All of the BMD matrices are truncated to omit empty modes, accounting for the fact that we based the construction of the BdG Fock states given in Eq. (5.218) on the truncated product states. The r and s matrix subscripts indicate a restriction of the column index set to the respective index set, r or s . For instance, U_r is a $d \times R$ matrix

$$U_r := \begin{pmatrix} U_{1m_1} & \dots & U_{1m_r} \\ \vdots & \ddots & \vdots \\ U_{dm_1} & \dots & U_{dm_r} \end{pmatrix}, \quad (5.220)$$

formed by the columns $m_1, \dots, m_r$ of U . The restriction to the specified index set is always applied before any other matrix operations, such as taking the transpose, adjoint or complex conjugate. The sign $\pm$ in Eq. (5.217) is given by

$$\pm = (-1)^{\frac{(o(o-1)+s(s-1))}{2}}, \quad (5.221)$$

where $o = d_{\text{F}} + 2d_{\text{P}}$ is the number of non-empty BdG modes in the vacuum $|0'\rangle_b^{\text{p}}$ and r is the number of excitations in $|n'_r\rangle_b$. The normalisation factors in Eq. (5.217) are given by

$$\mathcal{N} = \prod_{p=1}^{d_{\text{P}}} v_p \quad \text{and} \quad \mathcal{N}' = \prod_{p'=1'}^{d_{\text{P}'}} v_{p'}. \quad (5.222)$$

5.5 Majorana Modes in BdG Theories

In quantum field theory, Majorana fermions are particles associated with a special kind of quantum field. These Majorana type quantum fields are exotic in that they are real, which famously means that Majorana particles are their own antiparticles. In condensed matter theory, and in particular in BdG theories, the redundant structure of Nambu space allows us to rewrite BdG Hamiltonians in terms of certain quasi-fermionic operators that are self-adjoint and in this sense reminiscent of Majorana fermions. We encountered this before when we illustrated that the algebra of operators on the fermionic Fock space $\mathcal{F}(\mathcal{H})$ of an d -dimensional single-particle Hilbert space $\mathcal{H}$ resembles the Clifford algebra $\text{Cl}_{2d}(\mathbb{R})$. However, the special thing about BdG theories is not merely their formulation in Nambu space, which admits a Majorana basis, but their potential for Majorana *eigenstates*. To see this, consider a BdG Hamiltonian

$$H_{\text{BdG}} = \Psi^\dagger h_{\text{BdG}} \Psi = (\mathbf{c}^\dagger \mathbf{c}) \begin{pmatrix} T & \Delta^\dagger \\ \Delta & -T^* \end{pmatrix} \begin{pmatrix} \mathbf{c} \\ \mathbf{c}^\dagger \end{pmatrix}, \quad (5.223)$$

where $\Psi = (c_1 \dots c_d c_1^\dagger \dots c_d^\dagger)^\top$ denotes the Nambu spinor of the d elementary fermion annihilation and creation operators c_j and $c_j^\dagger$. A Bogoliubov diagonalisation

$$H_{\text{BdG}} = \Phi^\dagger E_{\text{BdG}} \Phi = (\mathbf{b}^\dagger \mathbf{b}) \begin{pmatrix} E & 0 \\ 0 & -E \end{pmatrix} \begin{pmatrix} \mathbf{b} \\ \mathbf{b}^\dagger \end{pmatrix} \quad (5.224)$$

of H_{BdG} yields $2d$ complex Bogoliubov eigenstates $b_1, \dots, b_d, b_1^\dagger, \dots, b_d^\dagger$, where the b_j and the $b_j^\dagger$ are understood as the quasiparticle and quasihole eigenstates, respectively. Now, we may unitarily combine particle-hole conjugate pairs of Bogoliubov state operators into new operators

$$\gamma_{(j,1)} = b_j^\dagger + b_j \quad \text{and} \quad \gamma_{(j,2)} = i(b_j^\dagger - b_j). \quad (5.225)$$

In this way, we can express the BdG problem that was originally stated in terms of d particle-hole conjugate pairs of complex Bogoliubov fermions through $2d$ operators that are self-adjoint

$$\gamma_{(j,1)}^\dagger = (b_j^\dagger + b_j)^\dagger = b_j + b_j^\dagger = \gamma_{(j,1)} \quad \text{and} \quad \gamma_{(j,2)}^\dagger = [i(b_j^\dagger - b_j)]^\dagger = -i(b_j - b_j^\dagger) = \gamma_{(j,2)} \quad (5.226)$$

and fulfil

$$\begin{aligned} \{\gamma_{(j,k)}, \gamma_{(l,m)}\} &= \{i^{(k-1)}(b_j^\dagger + (-1)^{(k-1)}b_j), i^{(m-1)}(b_l^\dagger + (-1)^{(m-1)}b_l)\} \\ &= i^{(k+m-2)} \left(\{b_j^\dagger, b_l^\dagger\} + (-1)^{(m-1)}\{b_j^\dagger, b_l\} \right. \\ &\quad \left. + (-1)^{(k-1)}\{b_j, b_l^\dagger\} + (-1)^{(k+m-2)}\{b_j, b_l\} \right) \\ &= i^{(k+m-2)} \left((-1)^{(m-1)}\delta_{jl} + (-1)^{(k-1)}\delta_{jl} \right) \\ &= i^{(k+m-2)} \left((-1)^{(m-1)} + (-1)^{(k-1)} \right) \delta_{jl} \\ &= \begin{cases} 2 & \text{for } j = l, k = m \\ 0 & \text{else,} \end{cases} \end{aligned} \quad (5.227)$$

where $j, l = 1, \dots, d$ and $k, m = 1, 2$. If we introduce multi-indices $\alpha = (j, k)$ and $\beta = (l, m)$, the Majorana algebra simplifies to the Clifford algebra of $2d$ generators

$$\{\gamma_\alpha, \gamma_\beta\} = 2\delta_{\alpha\beta}. \quad (5.228)$$

The self-adjointness of the γ operators is what earns them the name of Majorana operators. However, the anticommutation relations in Eq. (5.228) are not quite the same as the fermionic anticommutation relations: even though different γ operators anticommute properly as

$$\{\gamma_\alpha, \gamma_\beta\} = \{\gamma_\alpha^\dagger, \gamma_\beta\} = \{\gamma_\alpha, \gamma_\beta^\dagger\} = \{\gamma_\alpha^\dagger, \gamma_\beta^\dagger\} = 0, \quad (5.229)$$

for $\alpha \neq \beta$, the anticommutators for $\alpha = \beta$ are highly irregular. To begin with, the anticommutators

$$\{\gamma_\alpha^\dagger, \gamma_\alpha\} = \{\gamma_\alpha, \gamma_\alpha^\dagger\} = 2 \quad (5.230)$$

deviate from the canonic fermionic anticommutator relations by a factor of two. This factor could of course be removed by simply renormalising the γ operators by a factor of $1/\sqrt{2}$, but this is not usually done. We will soon see why. On top of that, the self-adjointness of the γ operators means that the formally sensible relations in Eq. (5.230) are also equivalent to

$$\{\gamma_\alpha, \gamma_\alpha\} = \{\gamma_\alpha^\dagger, \gamma_\alpha^\dagger\} = 2, \quad (5.231)$$

as indicated in Eq. (5.228). These anticommutator relations are not even formally fermionic anymore. They tell us that the individual γ operators fail to square to zero. Instead, they square to one since

$$\{\gamma_\alpha, \gamma_\alpha\} = \gamma_\alpha \gamma_\alpha + \gamma_\alpha \gamma_\alpha = 2\gamma_\alpha^2 = 2 \implies \gamma_\alpha^2 = 1. \quad (5.232)$$

Note that a renormalised version of the γ operators that satisfies $\{\gamma_\alpha, \gamma_\beta\} = \delta_{\alpha\beta}$ would square to one half instead. However, one half is an inconvenient value for γ_α^2 because it makes expressions like operator exponentials

$$e^{C\gamma_\alpha\gamma_\beta} \quad (5.233)$$

a fair bit harder to evaluate. For such purposes it is far more practical to choose the γ operators such that $\gamma_\alpha^2 \stackrel{!}{=} 1$, i.e. if $\{\gamma_\alpha, \gamma_\beta\} = 2\delta_{\alpha\beta}$ instead of $\{\gamma_\alpha, \gamma_\beta\} = \delta_{\alpha\beta}$. In a sense, we sacrifice a measure of resemblance to the fermionic algebra for better manageability. This is a favourable deal for us because said resemblance was only ever formal to begin with. The mere fact that the γ operators square to a non-zero constant presents a far more serious problem: it means that the γ number operator

$$n_\gamma := \gamma^\dagger \gamma = \gamma \gamma = 1 \quad (5.234)$$

is *constant*. As a result, there are no number states available for the γ operators and they cannot be accomodated as proper quasiparticle states in a fermionic Fock space. This is an important difference to the Majorana fermions from quantum field theory, which are of course real fermions.⁷

Now, the Majorana states usually fail to be eigenstates of the BdG Hamiltonian. To see this, we consider the matrix representation of the BdG eigenvalue equations for quasiparticle and quasihole Bogoliubov modes. These read

$$h_{\text{BdG}} \beta_j = +E_j \beta_j \quad \text{and} \quad h_{\text{BdG}} \eta_j = -E_j \eta_j, \quad (5.235)$$

where h_{BdG} denotes the BdG Hamiltonian matrix and where β_j and η_j denote the coefficient column vectors of the quasiparticle and quasihole Bogoliubov modes, i.e.

$$\beta_j = \begin{pmatrix} U_{1j} \\ \vdots \\ U_{dj} \\ V_{1j} \\ \vdots \\ V_{dj} \end{pmatrix} \quad \text{and} \quad \eta_j = \begin{pmatrix} V_{1j}^* \\ \vdots \\ V_{dj}^* \\ U_{1j}^* \\ \vdots \\ U_{dj}^* \end{pmatrix}, \quad (5.236)$$

with which

$$b_j = \beta_j^\dagger \Psi = \sum_{k=1}^d \left(U_{kj}^* c_k + V_{kj}^* c_k^\dagger \right) \quad \text{and} \quad b_j^\dagger = \eta_j^\dagger \Psi = \sum_{k=1}^d \left(V_{kj} c_k + U_{kj} c_k^\dagger \right). \quad (5.237)$$

In terms of the coefficient vectors, Eq. (5.225) can be written as

$$\gamma_{(j,1)} = (\eta_j^\dagger + \beta_j^\dagger) \Psi =: \gamma_{(j,1)}^\dagger \Psi \quad \text{and} \quad \gamma_{(j,2)} = i(\eta_j^\dagger - \beta_j^\dagger) \Psi =: \gamma_{(j,2)}^\dagger \Psi, \quad (5.238)$$

⁷In every sense of the word “real”.

which motivates the interpretation of the coefficient vectors

$$\boldsymbol{\gamma}_{(j,1)} = \boldsymbol{\beta}_j + \boldsymbol{\eta}_j \quad \text{and} \quad \boldsymbol{\gamma}_{(j,2)} = i(\boldsymbol{\beta}_j - \boldsymbol{\eta}_j) \quad (5.239)$$

as representatives of the Majorana modes. With these, we get

$$h_{\text{BdG}} \boldsymbol{\gamma}_{(j,1)} = h_{\text{BdG}}(\boldsymbol{\beta}_j + \boldsymbol{\eta}_j) \quad \text{and} \quad h_{\text{BdG}} \boldsymbol{\gamma}_{(j,2)} = h_{\text{BdG}} i(\boldsymbol{\beta}_j - \boldsymbol{\eta}_j) \quad (5.240)$$

$$= E_j(\boldsymbol{\beta}_j - \boldsymbol{\eta}_j) \quad \quad \quad = E_j i(\boldsymbol{\beta}_j + \boldsymbol{\eta}_j) \quad (5.241)$$

$$= -iE_j \boldsymbol{\gamma}_{(j,2)} \quad \quad \quad = iE_j \boldsymbol{\gamma}_{(j,1)} \quad (5.242)$$

so the pair of Majorana vectors $\boldsymbol{\gamma}_{(j,1)}$ and $\boldsymbol{\gamma}_{(j,2)}$ is swapped and equipped with an imaginary prefactor of $\pm iE_j$ under the action of h_{BdG} . These equations can only ever be understood as eigenvalue equations if $E_j = 0$, i.e. if the Majorana operators are made up of Bogoliubov quasiparticle and quasihole operators that belong to zero energy eigenstates. We will call Majorana modes that correspond to zero energy eigenstates *Majorana zero modes* (MZMs) and their operators MZM operators.

Note that zero energy Bogoliubov states are quite rare in BdG systems and appear only on special occasions. This is why the general narrative regarding Majorana modes in condensed matter systems is often summarised as follows: while it is always possible to rewrite a given BdG Hamiltonian in terms of Majorana operators, it is rare that these Majorana operators correspond to eigenstates of the Hamiltonian. That is to say, we can always formulate a BdG problem in terms of inseparably paired finite energy Majorana states, but we seldom encounter unpaired zero energy Majorana modes. Moreover, we note that in general there is no natural prescription on how to form complex fermion operators out of a given set of Majorana operators, so there is no natural way to rewrite a given Majorana Hamiltonian in the form of a BdG Hamiltonian.

The construction scheme for Majorana operators that we put forth in Eq. (5.225) is of course not unique. Since the combination of complex Bogoliubov operators into self-adjoint Majorana operators is usually a purely mathematical endeavour without immediate physical relevance, we are free to write down other, more involved schemes as long as the quasi-fermionic anticommutation relations remain intact. Of course, there is little to no reason to make things even more complicated for finite energy pairs of Majorana modes. However, this changes in the presence of complex *Bogoliubov* modes with zero energy. As we have seen, the MZMs that are constructed from such zero energy Bogoliubov modes *are* eigenstates of the Hamiltonian. Thus, the remaining degree of freedom in their construction becomes interesting. Let us consider a general BdG system with n complex Bogoliubov modes $b_1, \dots, b_n$ at zero energy. We want to find *all possible* sets of $2n$ MZM operators $\gamma_1, \dots, \gamma_{2n}$ that satisfy

$$\{\gamma_j, \gamma_k\} = 2\delta_{jk} \quad \text{and} \quad \gamma_j = \gamma_j^\dagger, \quad (5.243)$$

and are constructed from the n pairs of Bogoliubov creation and annihilation operators of the zero energy quasiparticle and quasihole Bogoliubov modes. To do this, we consider the most general ansatz for such operators,

$$\gamma_j = \sum_{k=1}^n (\alpha_{jk} b_k + \beta_{jk} b_k^\dagger), \quad (5.244)$$

with coefficients $\alpha_{jk}, \beta_{jk} \in \mathbb{C}$ for all $j, k = 1, \dots, n$. The self-adjointness requirement $\gamma_j = \gamma_j^\dagger$ immediately tells us that $\beta_{jk} \stackrel{!}{=} \alpha_{jk}^*$, so we get

$$\gamma_j = \sum_{k=1}^n (\alpha_{jk} b_k + \alpha_{jk}^* b_k^\dagger). \quad (5.245)$$

Additionally, we have to ensure the Clifford algebra requirement

$$\{\gamma_j, \gamma_l\} \stackrel{!}{=} 2\delta_{jl}, \quad (5.246)$$

which, in terms of Eq. (5.245), becomes

$$\begin{aligned}
\{\gamma_j, \gamma_l\} &= \left\{ \sum_{k=1}^n (\alpha_{jk} b_k + \alpha_{jk}^* b_k^\dagger), \sum_{m=1}^n (\alpha_{lm} b_m + \alpha_{lm}^* b_m^\dagger) \right\} \\
&= \sum_{k,m=1}^n \left(\alpha_{jk} \alpha_{lm} \{b_k, b_m\} + \alpha_{jk} \alpha_{lm}^* \{b_k, b_m^\dagger\} + \alpha_{jk}^* \alpha_{lm} \{b_k^\dagger, b_m\} + \alpha_{jk}^* \alpha_{lm}^* \{b_k^\dagger, b_m^\dagger\} \right) \\
&= \sum_{k,m=1}^n (\alpha_{jk} \alpha_{lm}^* + \alpha_{jk}^* \alpha_{lm}) \delta_{km} \\
&= \sum_{k=1}^n (\alpha_{jk} \alpha_{lk}^* + \alpha_{jk}^* \alpha_{lk}) \\
&= 2 \sum_{k=1}^n \text{Re}(\alpha_{jk} \alpha_{lk}^*) \\
&\stackrel{!}{=} 2\delta_{jl},
\end{aligned} \tag{5.247}$$

where $\text{Re}(\cdot)$ denotes the real part. These conditions can be reduced to

$$1 \stackrel{!}{=} \sum_{k=1}^n |\alpha_{jk}|^2 \quad \text{and} \quad 0 \stackrel{!}{=} \sum_{k=1}^n \text{Re}(\alpha_{jk} \alpha_{lk}^*) \tag{5.248}$$

for all $j \neq l \in \{1, \dots, n\}$. The second condition ensures that different MZM operators anticommute, while the first one guarantees that the individual MZM operators square to one. If we use these conditions in a system with $n = 1$ zero energy Bogoliubov mode b_1 , we get two MZMs

$$\gamma_1 = \alpha_{11} b_1 + \alpha_{11}^* b_1^\dagger \quad \text{and} \quad \gamma_2 = \alpha_{21} b_1 + \alpha_{21}^* b_1^\dagger \tag{5.249}$$

with constraints

$$|\alpha_{11}|^2 = |\alpha_{21}|^2 \stackrel{!}{=} 1 \quad \text{and} \quad \text{Re}(\alpha_{11} \alpha_{21}^*) \stackrel{!}{=} 0. \tag{5.250}$$

The first equation tells us that the α_j are phase factors

$$\alpha_{j1} = e^{i\phi_j}, \tag{5.251}$$

where $\phi_j \in \mathbb{R}$. If we plug this into the second condition, we get

$$\text{Re}(\alpha_{11} \alpha_{21}^*) \stackrel{!}{=} \text{Re}(e^{i(\phi_1 - \phi_2)}) \stackrel{!}{=} 0, \tag{5.252}$$

which tells us that the relative phase $\Delta\phi$ between α_{11} and α_{21} must be

$$\Delta\phi := \phi_1 - \phi_2 \stackrel{!}{=} \pm \frac{\pi}{2}. \tag{5.253}$$

With this, we can write

$$\phi \equiv \phi_1 \quad \text{and} \quad \phi_2 = \phi_1 - \Delta\phi = \phi \pm \frac{\pi}{2}, \tag{5.254}$$

and determine the general coefficient sets

$$(\alpha_{11}, \alpha_{11}^*) = (e^{i\phi}, e^{-i\phi}) \quad \text{and} \quad (\alpha_{21}, \alpha_{21}^*) = (e^{i(\phi \pm \pi/2)}, e^{-i(\phi \pm \pi/2)}). \tag{5.255}$$

Substituting these into Eq. (5.249) then yields the general form

$$\gamma_1 = e^{i\phi} b + e^{-i\phi} b^\dagger \quad \text{and} \quad \gamma_2 = e^{i(\phi \pm \pi/2)} b + e^{-i(\phi \pm \pi/2)} b^\dagger \tag{5.256}$$

of the two MZM operators in a system with one complex Bogoliubov zero mode. Of course, this can be simplified to

$$\gamma_1 = e^{i\phi} b + e^{-i\phi} b^\dagger \quad \text{and} \quad \gamma_2 = \pm i (e^{i\phi} b - e^{-i\phi} b^\dagger), \tag{5.257}$$

which, for $\phi = 0$, reproduces the canonic definition of Majorana operators given in Eq. (5.225). However, it is easy to come up with other valid Majorana operators like

$$\gamma_1 = e^{i\frac{\pi}{4}}b + e^{-i\frac{\pi}{4}}b^\dagger \quad \text{and} \quad \gamma_2 = e^{i\frac{3\pi}{4}}b + e^{-i\frac{3\pi}{4}}b^\dagger. \quad (5.258)$$

This particular version of Majorana coefficients is distinguished by the fact that it represents the most “symmetric” family of Majorana coefficients: if we draw

$$\mathbb{C} \supset \mathbb{S}_\mathbb{C}^1 := \{z \in \mathbb{C} : |z| = 1\} \quad (5.259)$$

into the complex plane and mark the positions of the four MZM coefficients on it, then the set

$$(\alpha_1, \alpha_1^*, \alpha_2, \alpha_2^*) = (e^{i\frac{\pi}{4}}, e^{-i\frac{\pi}{4}}, e^{i\frac{3\pi}{4}}, e^{-i\frac{3\pi}{4}}) \quad (5.260)$$

of coefficients is the only one that evenly partitions $\mathbb{S}_\mathbb{C}^1$ into quadrants. Later, we will find that this maximally symmetric coefficient family is the one that is usually preferred by numeric solvers in systems with a single zero energy Bogoliubov mode. Overall, the above construction scheme provides us with a continuous family of possible Majorana operators that are qualitatively the same, because they are constructed from the same two Bogoliubov operators b and $b^\dagger$. This changes when there is more than one zero energy Bogoliubov mode.

Let us consider a BdG system with $n = 2$ Bogoliubov modes b_1 and b_2 at zero energy. The most natural way to write down four MZM operators

$$\begin{aligned} \gamma_1 &= \alpha_{11}b_1 + \alpha_{11}^*b_1^\dagger + \alpha_{12}b_2 + \alpha_{12}^*b_2^\dagger, & \gamma_2 &= \alpha_{21}b_1 + \alpha_{21}^*b_1^\dagger + \alpha_{22}b_2 + \alpha_{22}^*b_2^\dagger \\ \gamma_3 &= \alpha_{31}b_1 + \alpha_{31}^*b_1^\dagger + \alpha_{32}b_2 + \alpha_{32}^*b_2^\dagger, & \gamma_4 &= \alpha_{41}b_1 + \alpha_{41}^*b_1^\dagger + \alpha_{42}b_2 + \alpha_{42}^*b_2^\dagger \end{aligned} \quad (5.261)$$

is to repeat the $n = 1$ construction for b_1 and b_2 individually, i.e. to choose

$$\begin{aligned} (\alpha_{11}, \alpha_{11}^*, \alpha_{12}, \alpha_{12}^*) &= (e^{i\phi_1}, e^{-i\phi_1}, 0, 0), & (\alpha_{21}, \alpha_{21}^*, \alpha_{22}, \alpha_{22}^*) &= (e^{i(\phi_1 \pm \pi/2)}, e^{-i(\phi_1 \pm \pi/2)}, 0, 0) \\ (\alpha_{31}, \alpha_{31}^*, \alpha_{32}, \alpha_{32}^*) &= (0, 0, e^{i\phi_2}, e^{-i\phi_2}), & (\alpha_{41}, \alpha_{41}^*, \alpha_{42}, \alpha_{42}^*) &= (0, 0, e^{i(\phi_2 \pm \pi/2)}, e^{-i(\phi_2 \pm \pi/2)}), \end{aligned} \quad (5.262)$$

such that

$$\begin{aligned} \gamma_1 &= e^{i\phi_1}b_1 + e^{-i\phi_1}b_1^\dagger, & \gamma_2 &= e^{i(\phi_1 \pm \pi/2)}b_1 + e^{-i(\phi_1 \pm \pi/2)}b_1^\dagger \\ \gamma_3 &= e^{i\phi_2}b_2 + e^{-i\phi_2}b_2^\dagger, & \gamma_4 &= e^{i(\phi_2 \pm \pi/2)}b_2 + e^{-i(\phi_2 \pm \pi/2)}b_2^\dagger. \end{aligned} \quad (5.263)$$

One can verify that the coefficients in Eq. (5.262) readily satisfy the conditions from Eq. (5.248), so that the naive definition of two “independent” pairs of MZMs does indeed produce valid Majorana operators. However, these are by far not the only Majorana constructions allowed. In particular, we could have based the independent Majorana pair scheme on unitarily rotated zero energy Bogoliubov operators

$$\begin{pmatrix} \tilde{b}_1 \\ \tilde{b}_2 \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} b_1 + b_2 \\ b_1 - b_2 \end{pmatrix}. \quad (5.264)$$

Due to the previous analysis, we know that the new Majorana operators

$$\begin{aligned} \tilde{\gamma}_1 &= \left(e^{i\phi_1}\tilde{b}_1 + e^{-i\phi_1}\tilde{b}_1^\dagger \right), & \tilde{\gamma}_2 &= \left(e^{i(\phi_1 \pm \pi/2)}\tilde{b}_1 + e^{-i(\phi_1 \pm \pi/2)}\tilde{b}_1^\dagger \right) \\ \tilde{\gamma}_3 &= \left(e^{i\phi_2}\tilde{b}_2 + e^{-i\phi_2}\tilde{b}_2^\dagger \right), & \tilde{\gamma}_4 &= \left(e^{i(\phi_2 \pm \pi/2)}\tilde{b}_2 + e^{-i(\phi_2 \pm \pi/2)}\tilde{b}_2^\dagger \right) \end{aligned} \quad (5.265)$$

are still well-defined.

6 – Anyons

In the vast landscape of quantum particles, there exists one particularly curious species known as anyons. A utilitarian definition of anyons is as follows: while fermions obey Fermi–Dirac exchange statistics and bosons obey Bose–Einstein exchange statistics, anyons obey *any* exchange statistics. Clearly, these kinds of particles are nowhere near as common in nature as fermions and bosons, otherwise they would be as familiar to everyone as fermions and bosons. As we will soon see, the reason for this is that anyons can only occur in two spatial dimensions and we just so happen to live in a world of three spatial dimensions. Thus, it is only under rather specific circumstances – when the motion in one spatial dimension is somehow “frozen out” – that anyons may emerge as quasiparticle excitations.

Consider a quantum state of some number of identical particles. The *exchange statistics* of the particles refers to the phase picked up by the state when two of the identical particles are exchanged. However, this definition is ambiguous. Does it refer to the phase acquired under a formal, instantaneous permutation of the particles, or does it refer to the phase that results when any two particles are adiabatically moved around each other such that they swap positions in the process? There is a philosophical debate to be had about which of these notions offers a more accurate description of nature. Fortunately, both interpretations turn out to be equivalent in three or more spatial dimensions and we do not have to concern ourselves with this distinction there. However, this is not the case in two dimensions. In the following, we will show how the concepts formal instantaneous and adiabatic real-space exchange statistics differ in two spatial dimensions. In particular, we will address the phases that result from adiabatic exchanges. This part is mainly based on Ref. [95].

The most natural way to understand statistics under adiabatic exchange of particles is to think about it in terms of path integrals. Recall that in the path integral formulation of quantum mechanics, the probability amplitude for a single particle to go from point $\mathbf{x}$ to point $\mathbf{y}$ is given by

$$A[\mathbf{x} \rightarrow \mathbf{y}] = \int_{\mathbf{x}}^{\mathbf{y}} \mathcal{D}\mathbf{r} e^{iS[\mathbf{r}, \dot{\mathbf{r}}]}, \quad (6.1)$$

where $\mathcal{D}\mathbf{r}$ denotes integration over all possible paths from $\mathbf{x}$ to $\mathbf{y}$, and where

$$S[\mathbf{r}, \dot{\mathbf{r}}] := \int dt \mathcal{L}[\mathbf{r}(t), \dot{\mathbf{r}}(t)] \quad (6.2)$$

is the action given in terms of the Lagrangian $\mathcal{L}[\mathbf{r}(t), \dot{\mathbf{r}}(t)]$ of the system. Note that the integration over all possible paths explicitly includes *discontinuous* paths. However, if the position $\mathbf{x}$ of a particle changes discontinuously, the position derivative $\dot{\mathbf{x}}$ and hence the Lagrangian $\mathcal{L}[\mathbf{r}(t), \dot{\mathbf{r}}(t)]$ become singular. As a result, the action $S[\mathbf{r}, \dot{\mathbf{r}}]$ along discontinuous paths is usually argued to be giant. A giant action leads to a rapidly rotating path contribution $e^{iS[\mathbf{r}, \dot{\mathbf{r}}]}$ in Eq. (6.1), which, in the spirit of the stationary phase approximation, is assumed to interfere destructively and can therefore be omitted. For this reason, we may limit the following considerations to continuous paths.

Now consider a system of N identical particles in an n spatial dimensions. For $N > 1$, the indistinguishability of the particles makes the construction of probability amplitudes a bit more subtle. Since permutations of identical particles produce physically equivalent states, the probability amplitude of a transition from a state with N particles at positions $(\mathbf{x}_1, \dots, \mathbf{x}_N)$ to a state with N particles at positions $(\mathbf{y}_1, \dots, \mathbf{y}_N)$ must account for permutations of the particles in the final state. Thus, we have

$$A[(\mathbf{x}_1, \dots, \mathbf{x}_N) \rightarrow (\mathbf{y}_1, \dots, \mathbf{y}_N)] = \sum_{\pi \in S_N} \int_{(\mathbf{x}_1, \dots, \mathbf{x}_N)}^{(\mathbf{y}_{\pi(1)}, \dots, \mathbf{y}_{\pi(N)})} \mathcal{D}(\mathbf{r}_1, \dots, \mathbf{r}_N) e^{iS[\mathbf{r}_1, \dot{\mathbf{r}}_1, \dots, \mathbf{r}_N, \dot{\mathbf{r}}_N]}, \quad (6.3)$$

where the extra sum extends over the elements π of the symmetric group S_N of N elements in order to include all possible permutations of end points. Before we use Eq. (6.3) to analyse the possible adiabatic exchange statistics of particles in n spatial dimensions, we impose one more constraint on the paths included in Eq. (6.3). By invoking a hard-core repulsion between the particles, we effectively exclude

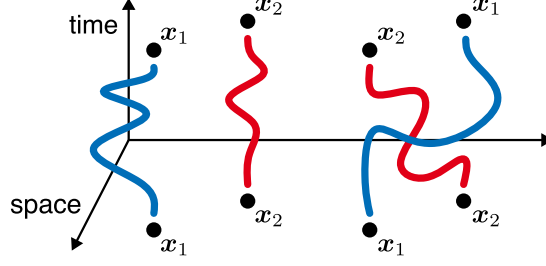


Figure 6.1: Exchangeless and single exchange paths in $(n + 1)$ dimensional spacetime. The n spatial dimensions are depicted as a plane.

intersecting paths along which any two particles can occupy the same region in space. This additional restriction can be justified in several ways. Physically, pairs of identical particles often experience a strong repulsive interaction that keeps them separated. Philosophically, one could argue that the very concept of adiabatic exchange statistics essentially requires the particles to remain well-separated, as otherwise their real-space exchange cannot be tracked in a meaningful way. Note that by including the hard-core constraint, we exclude cases with $n = 1$ spatial dimension from the discussion, since an adiabatic exchange of hard-core particles is impossible there. Accepting this caveat, the probability amplitude for an adiabatic exchange of $N = 2$ identical particles in $n \geq 2$ spatial dimensions becomes

$$A[(\mathbf{x}_1, \mathbf{x}_2) \rightarrow (\mathbf{x}_2, \mathbf{x}_1)] = \int_{(\mathbf{x}_1, \mathbf{x}_2)}^{(\mathbf{x}_1, \mathbf{x}_2)} \mathcal{D}(\mathbf{r}_1, \mathbf{r}_2) e^{i\mathcal{S}[\mathbf{r}_1, \dot{\mathbf{r}}_1, \mathbf{r}_2, \dot{\mathbf{r}}_2]} + \int_{(\mathbf{x}_1, \mathbf{x}_2)}^{(\mathbf{x}_2, \mathbf{x}_1)} \mathcal{D}(\mathbf{r}_1, \mathbf{r}_2) e^{i\mathcal{S}[\mathbf{r}_1, \dot{\mathbf{r}}_1, \mathbf{r}_2, \dot{\mathbf{r}}_2]}, \quad (6.4)$$

where the path integral takes into account all continuous, non-intersecting paths from $(\mathbf{x}_1, \mathbf{x}_2)$ to $(\mathbf{x}_2, \mathbf{x}_1)$ and $(\mathbf{x}_1, \mathbf{x}_2)$. Two simple examples of such paths are sketched in Fig. 6.1. The key observation is that the exchange statistics of identical particles should not depend on the geometric details of exchange paths: if two exchange paths can be continuously deformed into one another, they should give rise to the same *statistical* phases. This suggests that the possible exchange statistics of identical particles in n spatial dimensions are somehow related to the *homotopically distinct* exchange paths, giving their classification a distinctly topological flavour. In fact, the challenge of counting distinct exchange paths is closely reminiscent of the problem that led to the definition of the fundamental group. If we consider the *relative* position $\mathbf{x} = \mathbf{x}_2 - \mathbf{x}_1$ of the two particles, then every path $\gamma : (\mathbf{x}_1, \mathbf{x}_2) \rightarrow (\mathbf{x}_1, \mathbf{x}_2)$ that returns both particles to their initial positions naturally defines a *closed* path $\Gamma : \mathbf{x} \rightarrow \mathbf{x}$ in the configuration space

$$C_2^n \equiv \mathbb{R}^n \setminus \{\mathbf{0}\} \quad (6.5)$$

of the relative position $\mathbf{x}$ in n spatial dimensions. Note, that we remove the origin from $\mathbb{R}^n$ because the hard-core constraint means that paths cannot intersect. Every path $\gamma : (\mathbf{x}_1, \mathbf{x}_2) \rightarrow (\mathbf{x}_1, \mathbf{x}_2)$ can therefore be naturally associated with an element $[\Gamma : \mathbf{x} \rightarrow \mathbf{x}]$ of the fundamental group $\pi_1(C_2^n)$. Unfortunately, the second class of paths, $\lambda : (\mathbf{x}_1, \mathbf{x}_2) \rightarrow (\mathbf{x}_2, \mathbf{x}_1)$, where the particles end up in each other's initial position, correspond to *open* paths $\Lambda : \mathbf{x} \rightarrow -\mathbf{x}$ in C_2^n that are not directly classified by $\pi_1(C_2^n)$. However, the indistinguishability of the particles tells us that all physical states characterised by $\mathbf{x}$ and $-\mathbf{x}$ are actually equivalent. To obtain the configuration space $\mathcal{C}_2^n$ of the relative position $\mathbf{x}$ of two *identical* particles in n spatial dimensions, we must therefore *identify* the relative positions $\mathbf{x}$ and $-\mathbf{x}$ in C_2^n , yielding

$$\mathcal{C}_2^n \equiv (\mathbb{R}^n \setminus \{\mathbf{0}\}) / \mathbb{Z}_2, \quad (6.6)$$

i.e. the punctured n -dimensional Euclidean space $C_2^n \equiv \mathbb{R}^n \setminus \{\mathbf{0}\}$ modulo the group action of $\mathbb{Z}_2$ that identifies $\mathbf{x} \simeq -\mathbf{x}$ for all $\mathbf{x} \in C_2^n$. Now all paths considered in Eq. (6.4) define closed paths in $\mathcal{C}_2^n$. Accordingly, every path also belongs to some equivalence class of the fundamental group $\pi_1(\mathcal{C}_2^n)$. This indicates that the possible statistics of identical particles in n spatial dimensions are determined by the fundamental group $\pi_1(\mathcal{C}_2^n)$. We can compute these fundamental groups by noting that the configuration spaces $\mathcal{C}_2^n$ in Eq. (6.6) are homotopy equivalent to the real projective spaces $\mathbb{RP}^{n-1}$ of one lower dimension, i.e.

$$\mathcal{C}_2^n \simeq \mathbb{RP}^{n-1}. \quad (6.7)$$

This follows from the fact that $\mathbb{R}^n \setminus \{\mathbf{0}\}$ deformation-retracts to the sphere $\mathbb{S}^{n-1}$ and the $\mathbb{Z}_2$ quotient of $\mathbb{S}^{n-1}$ immediately yields

$$\mathbb{S}^{n-1}/\mathbb{Z}_2 \cong \mathbb{RP}^{n-1} \quad (6.8)$$

by definition. Note that Eq.(6.7) expresses a homotopy equivalence, indicated by the symbol $\simeq$, whereas Eq.(6.8) establishes a homeomorphism, denoted by $\cong$. Equation (6.8) shows that $\mathbb{S}^{n-1}$ constitutes a double cover

$$\begin{aligned} \pi : \mathbb{S}^{n-1} &\rightarrow \mathbb{RP}^{n-1} \\ x &\mapsto [x] \end{aligned} \quad (6.9)$$

of $\mathbb{RP}^{n-1}$, where $[x]$ denotes the equivalence class $[x] = \{x, -x\}$ of antipodal points. In the language of fibre bundles, Eq. (6.9) defines a fibre bundle with base manifold $B = \mathbb{RP}^{n-1}$ and discrete fiber $F = \{x, -x\}$.¹ For $n \geq 3$, Eq. (6.9) can be used to determine the fundamental groups $\pi_1(\mathcal{C}_2^n) = \pi_1(\mathbb{RP}^{n-1})$ via the following theorem.

Theorem 6.0.1. Deck Group Isomorphism. *Let $\pi : D \rightarrow X$ be a double cover. The **deck group** G is the group of homeomorphisms $d : D \rightarrow D$ that preserve the covering structure, satisfying*

$$\pi \circ d = \pi. \quad (6.10)$$

If the double cover is universal (simply connected), the deck group G is isomorphic to the fundamental group $\pi_1(X)$ of the base space X , i.e.

$$G \simeq \pi_1(X). \quad (6.11)$$

Intuitively speaking, the deck transformations of a covering space permute the discrete fibres in a way that is compatible with the bundle structure. Since the covering space $\mathbb{S}^{n-1}$ of $\mathcal{C}_2^n$ is simply connected for $n \geq 3$, Thm. 6.0.1 tells us that

$$\pi_1(\mathcal{C}_2^n) \equiv \pi_1(\mathbb{RP}^{n-1}) \simeq G_{n-1} \quad \text{for } n \geq 3, \quad (6.12)$$

where G_{n-1} denotes the deck group of $\pi : \mathbb{S}^{n-1} \rightarrow \mathbb{RP}^{n-1}$ as a discrete fibre bundle. In that case, the deck groups G_n are directly determined by the group $\mathbb{Z}_2$ that defines the covering in Eq. (6.9) in the first place. Indeed, the only two homeomorphisms $d : \mathbb{S}^n \rightarrow \mathbb{S}^n$ fulfilling Eq. (6.10) for the double cover Eq. (6.9) are the identity map $\text{id} : \mathbb{S}^{n-1} \rightarrow \mathbb{S}^{n-1}$, $x \mapsto x$ and the antipodal map $a : \mathbb{S}^{n-1} \rightarrow \mathbb{S}^{n-1}$, $x \mapsto -x$, so that we have

$$G_{n-1} = \{\text{id}, a\} \simeq \mathbb{Z}_2, \quad (6.13)$$

and hence

$$\pi_1(\mathcal{C}_2^n) \equiv \pi_1(\mathbb{RP}^{n-1}) \simeq \mathbb{Z}_2 \quad (6.14)$$

for all $n \geq 3$. For $n = 2$, we cannot apply Thm. 6.0.1 as $\mathbb{S}^1$ is not simply connected and therefore not a universal cover of $\mathcal{C}_2^2 \equiv \mathbb{RP}^1$. The fundamental group $\pi_1(\mathbb{RP}^1)$ is still easy to find because the real projective line $\mathbb{RP}^1$ is *itself* homeomorphic to the circle $\mathbb{S}^1$, i.e. $\mathbb{RP}^1 \simeq \mathbb{S}^1$, such that

$$\pi_1(\mathbb{RP}^1) \simeq \pi_1(\mathbb{S}^1) \simeq \mathbb{Z}. \quad (6.15)$$

Combined, we therefore get

$$\pi_1(\mathcal{C}_2^n) \equiv \pi_1(\mathbb{RP}^{n-1}) = \begin{cases} \mathbb{Z} & \text{for } n = 2 \\ \mathbb{Z}_2 & \text{for } n \geq 3. \end{cases} \quad (6.16)$$

¹An analogous notion is true for all covering spaces.

In particular, we have $\pi_1(\mathcal{C}_2^3) \simeq \mathbb{Z}_2$ in $n = 3$ spatial dimensions, telling us that there are just two classes of homotopically distinct closed paths in $\mathcal{C}_2^3$, namely the direct and the single exchange paths. This means that every higher order exchange path can be continuously deformed into one of the two elementary paths shown in Fig. 6.1. For example, a double exchange of particles in $(3 + 1)$ dimensions is homotopically equivalent to the direct path: picturing a double exchange path in the style of Fig. 6.1, one can imagine sliding the first particle's doubly intertwined world line off the world line of the second particle. The probability amplitude in Eq. (6.3) for particles in $(3 + 1)$ dimensions can then be written as

$$A[(\mathbf{x}_1, \mathbf{x}_2) \rightarrow (\mathbf{x}_2, \mathbf{x}_1)] = \oint_0 \mathcal{D}\mathbf{R} e^{iS_0[\mathbf{R}, \dot{\mathbf{R}}]} + e^{i\phi} \oint_1 \mathcal{D}\mathbf{R} e^{iS_1[\mathbf{R}, \dot{\mathbf{R}}]}, \quad (6.17)$$

where we defined $\mathbf{R} := (\mathbf{r}_1, \mathbf{r}_2)$ and $\dot{\mathbf{R}} := (\dot{\mathbf{r}}_1, \dot{\mathbf{r}}_2)$ for convenience and wrote $\oint_n$ for the path integrals over the two classes of direct and single exchange paths. We also explicitly included a relative phase factor $e^{i\phi}$ that accounts for a possible topological phase difference between the two types of exchange paths. As was mentioned before, the statistical properties of the identical particles under adiabatic exchange must not depend on the geometrical details of the exchange paths, so only the relative phase factor $e^{i\phi}$ can encode the possible exchange statistics. The fact that a double exchange path of particles is homotopically equivalent to the direct path then yields the constraint

$$e^{2i\phi} \stackrel{!}{=} 1, \quad (6.18)$$

which implies that $\phi \in \{0, \pi\}$ giving rise to the familiar Bose–Einstein and Fermi–Dirac statistics available for identical particles in $(3 + 1)$ dimensional spacetime. This analysis carries over to all spatial dimensions of $n \geq 3$, showing that the only possible statistics in $n \geq 3$ spatial dimensions are of the Bose–Einstein and Fermi–Dirac types.

The special thing about $n = 2$ spatial dimensions is that $\pi_1(\mathcal{C}_2^2) \simeq \mathbb{Z}$. This means that there is an infinite number of homotopically distinct closed paths in $\mathcal{C}_2^2 \simeq \mathbb{RP}^1$. Every exchange path is labelled by an integer N indicating the number and direction (sign) of exchanges it involves, and any two exchange paths with $N \neq M$ are homotopically distinct. The probability amplitude in Eq. (6.3) for particles in $(2 + 1)$ dimensions can therefore be written as

$$A[(\mathbf{x}_1, \mathbf{x}_2) \rightarrow (\mathbf{x}_2, \mathbf{x}_1)] = \sum_{n \in \mathbb{Z}} e^{i\phi n} \oint_n \mathcal{D}\mathbf{R} e^{iS_n[\mathbf{R}, \dot{\mathbf{R}}]}, \quad (6.19)$$

where we introduced a sum over the classes of distinct exchange paths labelled by $n \in \mathbb{Z} \simeq \pi_1(\mathcal{C}_2^2)$ and wrote $\oint_n$ for the path integrals over the n -th class of exchange paths. Unlike before, Eq. (6.19) imposes no constraint on the relative phase factor $e^{i\phi}$ between equivalence classes of exchange paths: larger numbers of particle exchanges are no longer topologically required to yield the same statistical phase factor as the direct path. While $\phi \in \{0, \pi\}$ are still valid choices – it is possible to have bosons and fermions in two spatial dimensions – these are by far not the only choices allowed. In principle, *any* real phase $\phi \in \mathbb{R}$ is possible. Accordingly, we conclude that particles in $n = 2$ spatial dimensions may have *any* statistics, and refer to them *anyons*.

Observe that the fundamental group $\pi_1(\mathcal{C}_2^n) \simeq \mathbb{Z}_2$ determining the adiabatic exchange statistics of two identical particles in $n \geq 3$ spatial dimensions is the same as the symmetric group $S_2 \simeq \mathbb{Z}_2$ describing their formal permutation statistics, i.e.

$$\pi(\mathcal{C}_2^n) \simeq \mathbb{Z}_2 \simeq S_2 \quad \text{for } n \geq 3. \quad (6.20)$$

In $n = 2$ spatial dimensions, on the other hand, the adiabatic exchange statistics of identical particles is determined by $\pi_1(\mathcal{C}_2^2) \simeq \mathbb{Z}$. Thus, they *differ* from the formal permutation statistics described by $S_2 \simeq \mathbb{Z}_2$, and we have

$$\pi(\mathcal{C}_2^2) \simeq \mathbb{Z} \not\simeq \mathbb{Z}_2 \simeq S_2. \quad (6.21)$$

The above discussion is limited to the minimum number of $N = 2$ identical particles required for the definition of quantum statistics. For larger numbers of N identical particles, an analysis of the independent relative position vectors is less intuitive. Instead, one defines the so-called ordered configuration space

$$\text{Conf}_N(\mathbb{R}^n) := \left\{ (\mathbf{x}_1, \dots, \mathbf{x}_N) \in (\mathbb{R}^n)^N \mid \mathbf{x}_i \neq \mathbf{x}_j \text{ for all } i \neq j \right\}, \quad (6.22)$$

which contains all ordered N -tuples of pairwise distinct points (particle positions) in $\mathbb{R}^n$. To account for the indistinguishability of particles, we introduce an equivalence relation

$$(\mathbf{x}_1, \dots, \mathbf{x}_N) \sim (\mathbf{y}_1, \dots, \mathbf{y}_N) \quad \text{if} \quad (\mathbf{x}_1, \dots, \mathbf{x}_N) = (\mathbf{y}_{\pi(1)}, \dots, \mathbf{y}_{\pi(N)}), \quad (6.23)$$

for some element $\pi \in S_N$ of the symmetric (permutation) group S_N of N elements. Note that

$$\pi \cdot (\mathbf{x}_1, \dots, \mathbf{x}_N) = (\mathbf{x}_{\pi(1)}, \dots, \mathbf{x}_{\pi(N)}), \quad (6.24)$$

defines the natural action of S_N on $\text{Conf}_N(\mathbb{R}^n)$. The equivalence relation in Eq. (6.23) removes the ordering of the N -tuples by identifying all ordered N -tuples that contain the same N particle positions. The quotient of the ordered configuration space by the symmetric group then yields the so-called unordered configuration space

$$\text{UConf}_N(\mathbb{R}^n) := \text{Conf}_N(\mathbb{R}^n)/S_N. \quad (6.25)$$

For $N = 2$, $\text{Conf}_N(\mathbb{R}^n)$ is homotopy equivalent to C_2^n from Eq. (6.5), while $\text{UConf}_N(\mathbb{R}^n)$ is homotopy equivalent to C_2^n from Eq. (6.6) – in this sense, Eqs. (6.22) and (6.25) generalise Eqs. (6.5) and (6.6). Any adiabatic exchange among the N identical particles corresponds to a continuous path

$$(\mathbf{x}_1, \dots, \mathbf{x}_N) = (\mathbf{x}_{\pi(1)}, \dots, \mathbf{x}_{\pi(N)}), \quad (6.26)$$

which readily defines a closed path in $\text{UConf}_N(\mathbb{R}^n)$. The exchange statistics of N identical particles in n spatial dimensions is therefore captured by the fundamental group

$$\pi_1(\text{UConf}_N(\mathbb{R}^n)) =: B_N(\mathbb{R}^n), \quad (6.27)$$

which is called the N -strand braid group $B_N(\mathbb{R}^n)$ of $\mathbb{R}^n$ in the mathematical literature. Analogous to the previous discussion, Eq. (6.25) makes $\text{Conf}_N(\mathbb{R}^n)$ a covering space of $\text{UConf}_N(\mathbb{R}^n)$ but this time with deck group $G = S_N$. Moreover, the ordered configuration space $\text{Conf}_N(\mathbb{R}^n)$ turns out to be simply connected for $n \geq 3$, so that the covering becomes universal and we have

$$\pi_1(\text{UConf}_N(\mathbb{R}^n)) = B_N(\mathbb{R}^n) \simeq S_N \quad \text{for} \quad n \geq 3. \quad (6.28)$$

For $n = 2$, we instead find

$$\pi_1(\text{UConf}_N(\mathbb{R}^2)) = B_N(\mathbb{R}^2) \simeq B_N, \quad (6.29)$$

where B_N is the Artin braid group of N strands. Combined we get

$$\pi_1(\text{UConf}_N(\mathbb{R}^n)) = B_N(\mathbb{R}^n) = \begin{cases} B_N & \text{for } n = 2 \\ S_N & \text{for } n \geq 3, \end{cases} \quad (6.30)$$

which, noting that $B_2 \simeq \mathbb{Z}$ and $S_2 \simeq \mathbb{Z}_2$, readily generalises Eq. (6.16). This result tells us that in $n \geq 3$ spatial dimensions both the adiabatic exchange statistics and the formal permutation statistics are governed by the symmetric group and are therefore equivalent. In $(2 + 1)$ dimensions, the adiabatic exchange statistics are determined by the braid group and the equivalence with the formal permutation statistics is lost.

Some comments are in order. In the mathematical literature, the study of configuration spaces and their fundamental groups is not limited to $\mathbb{R}^n$ – instead, one considers arbitrary topological spaces X . As part of this generalisation, the term *N -strand braid group of X* was introduced to mean the fundamental group $B_N(X) \equiv \pi_1(\text{UConf}_N(X))$. Historically, the Artin braid group $B_N \equiv B_N(\mathbb{R}^2)$ from before was the first N -strand braid group to be studied, although the definition that Artin gave was much more rooted in algebra than in topology [96, 97]. In fact, the visual representation of $B_N(\mathbb{R}^2)$ is what inspired the terminology of “braids” in the first place. For these reasons, $B_N \equiv B_N(\mathbb{R}^2)$ is still known as *the* braid group today and caution is advised when navigating the literature on this subject.

6.1 Anyons and Spin

One of the more powerful results of quantum field theory is that it provides a framework to understand not only the emergence of spin but also its connection to quantum statistics: the spin-statistics theorem states that fields with half-integer spin must be quantised as fermions, while fields with integer spin must be quantised as bosons. So how does the spin-statistics theorem deal with anyonic statistics? In two spatial dimensions there is only one generator of angular momentum J_z – the projection of angular momentum onto the “missing” z -axis. An immediate consequence of this is that there is no helicity in $(2+1)$ dimensions, as the projection of S_z onto the xy -plane is always trivial. Since helicity is the only meaningful notion of spin for massless particles, it follows that massless particles in $(2+1)$ dimensions must be spinless. For particles with invariant mass M , spin can be unambiguously defined as

$$S = MS_z, \quad (6.31)$$

where S_z is the spin part of J_z . However, since there is only one generator of angular momentum, there is no non-Abelian Lie algebra that can restrict its possible eigenvalues $s \in \mathbb{R}$ and the spin of a massive particle in $(2+1)$ dimensions can in principle be anything. A particle ψ with any spin $s \in \mathbb{R}$ then transforms as

$$\psi \mapsto e^{i2\pi s} \psi, \quad (6.32)$$

under full 2π rotations. This recovers bosons for $s \in \mathbb{Z}$ and fermions for $s \in \mathbb{Z}/2$, but yields a new type of particle for any other $s \in \mathbb{R}$. This is awfully similar to the conclusion that particles in $(2+1)$ dimensions can have any statistics, and for good reason. In fact, there is another way to arrive at the conclusion that anyons may exist in two spatial dimensions. It starts from the observation that particles in $(2+1)$ dimensions can have *any* spin and concludes that they must be allowed to obey any statistics by invoking the spin-statistics theorem. The beauty of this argument is that it appreciates the weight of the spin-statistics theorem: a $(2+1)$ dimensional particle with some non-(half-)integer spin $s \in \mathbb{R}$ has to obey an accordingly exotic “continuous” statistics in order to remain in line with the spin-statistics theorem. In this sense, anyons smoothly interpolate between bosons and fermions both in the spin and the quantum statistical sense. For more Details, see, for instance, Refs. [98, 99].

6.2 Topological Quantum Computation with Anyons

Quantum computation is a method of information processing that utilises the superposition principle of quantum mechanics. The basic idea is simple: take the fundamental information unit of a classical computer, a bit which can be either 1 or 0, and turn it into a quantum object, so that it can exist not only as 1 or 0 but also as any superposition of these. The result is called a quantum bit, or qubit for short. In principle, every two-level system of states $|0\rangle$ and $|1\rangle$ that allows for coherent superpositions

$$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle, \quad (6.33)$$

with $\alpha, \beta \in \mathbb{C}$ can be used to represent a qubit. Here, $|\psi\rangle$ is called the qubit state and the two states $|0\rangle$ and $|1\rangle$ are called the logical states. Coherence means that the relative phase between the logical states $|0\rangle$ and $|1\rangle$ is well-defined and stable over time. It is essential because it enables the constructive and destructive interference necessary for quantum algorithms. Without coherence, the qubit loses its quantum behavior, and quantum computations become meaningless. In practice, coherence is a fragile property and is often compromised by noise or interactions with the environment. A logical gate on a qubit is implemented by a unitary operator $U \in \text{U}(2)$ that acts on the qubit state $|\psi\rangle$. For example, a NOT-gate U_{NOT} swaps the two logical states $|0\rangle$ and $|1\rangle$ in $|\psi\rangle$ as

$$U_{\text{NOT}}|0\rangle = |1\rangle \quad \text{and} \quad U_{\text{NOT}}|1\rangle = |0\rangle. \quad (6.34)$$

In this sense, it implements the quantum analogue of a classical bit flip operation. In the basis of the logical states, a generic qubit state $|\psi\rangle$ as in Eq. (6.33) takes the form $|\psi\rangle = (\alpha, \beta)^\top$ and U_{NOT} is

represented by the x -Pauli matrix $\mathcal{U}_{\text{NOT}} = \sigma_x$, giving

$$\mathcal{U}_{\text{NOT}} |\psi\rangle = \mathcal{U}_{\text{NOT}} \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = \begin{pmatrix} \beta \\ \alpha \end{pmatrix} \equiv |\bar{\psi}\rangle, \quad (6.35)$$

where $|\bar{\psi}\rangle$ denotes the “flipped” qubit state. The NOT-gate is also called an X -gate for this reason.

Topological quantum computation is an approach to quantum computing that is based on the exotic statistics of anyons. It builds on the observation that in some quantum systems, the presence of anyons leads to a degenerate coherent subspace $\mathcal{H}_0$, in which the state can only be evolved by adiabatically moving the anyons around each other. During such a sequence of adiabatic anyon exchanges, the anyon world lines tie up into so-called braids, that cannot be untangled due to the exotic statistical properties of the anyons. When a qubit is encoded in the degenerate coherent subspace $\mathcal{H}_0$ of the anyons, their adiabatic exchange can be used to implement quantum gates on $\mathcal{H}_0$ and the topology of the resulting anyonic world line knots provides powerful protection against local perturbations.

In the following, we give a brief overview over the theory behind topological quantum computation. To this end, we first introduce the algebraic theory of anyons, which provides a formal basis for the description of anyons and the unitary transformations realised by their braiding. After that, we discuss the definition and braiding of anyonic quasiparticles in the physical model of a topological p -wave superconductor.

6.3 Algebraic Theory of Anyons

Formally, a model of anyons is determined by a pair (S, T) of matrices called its modular data. These matrices encode essential physical properties of anyonic systems. The diagonal twist matrix T captures the so-called topological spin of the individual particles, specifying their transformation behaviour under 2π rotations. The elements $T_{ij} = \delta_{ij}\theta_j$ must be roots of unity, i.e. $\theta_j = \exp(i2\pi k/N_j)$ for $k, N_j \in \mathbb{Z}$, such that a 2π rotation of the j -th anyon type ψ_j is described by $\psi_j \mapsto \theta_j \psi_j$, cf. Eq. (6.32). The S matrix determines the quantum statistics of the particles, detailing the braiding between them. Together, the S and T matrices can be used to derive a set of anyon-combination rules known as *fusion rules*. Note that the converse is not true; there exist scenarios, such as in the well-known Ising and Fibonacci models, where the same set of fusion rules can correspond to multiple S and T matrices. This complexity gives rise to the concept of modular tensor categories (MTCs) and it is MTCs that provide a complete description of the algebraic properties for a given model of anyons [100]. Nonetheless, discussions of anyon models typically begin by stating the fusion rules in practice. This is what we will do in the following.

In condensed matter physics, anyons emerge as quasiparticles in some two-dimensional systems. There are two types of such anyon excitations: Abelian and non-Abelian anyons. In the much simpler case of Abelian anyons, the physical ground state is non-degenerate and an adiabatic exchange of anyonic quasiparticles can only furnish it with a $U(1)$ phase $\exp(i\phi)$. Abelian anyons are then called Abelian because $U(1)$ is an Abelian group. The much more interesting and challenging case is that of non-Abelian anyons, where the ground state is d -fold degenerate and we are dealing with a d -dimensional subspace $\mathcal{H}_0$ of ground states. The adiabatic exchange of anyonic quasiparticles will then induce a unitary $U(d)$ transformation on $\mathcal{H}_0$ that is no longer Abelian. For this part we closely follow Refs. [95, 101–103].

To describe a system of anyons, we first state the types of anyons in the system. We will represent the anyon types as $a = \{x_i\}_{i=0}^{N-1}$ and use $A = \{X_i\}_{i=0}^{N-1}$ to denote a representative set of anyons; the type of an anyon representative X_i is then x_i . Every anyonic system has a trivial anyon type that is usually denoted by $\mathbf{1}$. It represents the ground or vacuum state(s) of the system. In the list of anyon types above, we always set $x_0 = \mathbf{1}$. The fusion of two anyons is a process that is similar to the combination of two quantum spins to form a new total spin. In formula, we denote the fusion of two anyons X_1 and X_2 by

$$X_1 \otimes X_2. \quad (6.36)$$

Here, the $\otimes$ operation is both commutative and associative. Generally, the fusion of two anyons X_1 and X_2 with types x_1 and x_2 does not result in an anyon of a single, well-defined type x_i . Instead, the resulting anyons may be of several anyon types each with certain probabilities determined by the

aforementioned fusion rules. If the fusion of an anyon X_1 with every other anyon X_2 (including X_1 itself) always produces an anyon of the same type, then X_1 is called an Abelian anyon. The trivial particle $\mathbf{1}$ is Abelian as its fusion with any other particle X does not change the type of X , i.e. $\mathbf{1} \otimes x = x$ for every type x . If X_1 and X_2 are not Abelian, their fusion will produce anyons of more than one type and we say that the fusion has multiple fusion channels. We formally write the fusion result of any pair of anyons X_i and X_j as

$$X_i \otimes X_j = \bigoplus_{k=0}^{N-1} M_{ij}^k X_k. \quad (6.37)$$

The M_{ij}^k are non-negative integers called the fusion rules of the anyonic system. The fusion result of any pair of anyon types x_i, x_j is analogous to Eq. (6.37). Now, the coefficients in Eq. (6.37) that determine the fusion rules of an anyonic system can in principle be any non-negative integer. However, for most physical realisations of anyonic systems we have $M_{ij}^k \in \{0, 1\}$. If $M_{ij}^k = 0$ then fusing X_i with X_j can never yield X_k . If for all $X_i, X_j \in A$ there is only one M_{ij}^k that is different from zero, then the fusion outcome of each pair of anyons is unique and the model is Abelian. On the other hand, if for some pair X_i and X_j of anyons there are two or more fusion coefficients $M_{ij}^k \neq 0$ then the model is non-Abelian. Every anyon particle X_i has an antiparticle X_i^* defined by

$$X_i \otimes X_i^* = \mathbf{1}. \quad (6.38)$$

Note that an anyon particle X_i of type x_i may have an antiparticle X_i^* of a different anyon type $x_j \neq x_i$. If X_i and X_i^* are of the same type x_i , we call X_i self-dual. It is possible to view the integer coefficients M_{ij}^k as matrix entries $(M_i)_{jk}$ of a matrix with row and column indices j and k . The largest eigenvalue of this matrix is called the quantum dimension d_{x_i} of the anyon type x_i .

The term “quantum dimension” is a bit suspicious, so we will elaborate on it briefly. The key characteristic of non-Abelian anyons is that the fusion channels imply the existence of a subspace of degenerate ground states spanned by the different fusion outcomes. Consider a system with $N = 2$ anyons X_i and X_j of the same anyon type $x_i = x_j \equiv x$. Say X_i and X_j can fuse to several $X_k \in A$. Then we can formally define fully fused orthonormal states

$$|(X_i X_j); X_k\rangle \quad \text{with} \quad \langle (X_i X_j); X_{k_a} | (X_i X_j); X_{k_b} \rangle = \delta_{k_a k_b}, \quad (6.39)$$

where $(X_i X_j) := X_i \otimes X_j$ is a shorthand notation for the fusion of X_i and X_j and where k_a, k_b label all the possible fusion outcomes.² If there are d distinct fusion channels, then the system exhibits a d -dimensional subspace $\mathcal{H}_0[X_1, X_2]$ of degenerate ground states, which is spanned by the fully fused states. This subspace is accordingly called the fusion space. The fusion-space dimension typically grows when the number of non-Abelian anyons is increased. Indeed, in a system with N non-Abelian anyons $X_{i_1}, \dots, X_{i_N}$ of type x , the dimension of the fusion space $\mathcal{H}_0[X_{i_1}, \dots, X_{i_N}]$ roughly scales as

$$\dim(\mathcal{H}_0[X_{i_1}, \dots, X_{i_N}]) \xrightarrow{N \rightarrow \infty} d_x^N, \quad (6.40)$$

where d_x is a number that depends on the anyon type x . This scaling law is similar to the dimensional scaling of a tensor product of d_x -dimensional Hilbert spaces. For this reason, d_x is called the quantum dimension of the anyon type x . Conceptually, the quantum dimension is the asymptotic degeneracy per anyon of type x . For Abelian anyons the dimension of the space of ground states is equal to one no matter how many anyons there are in the system. Thus, all Abelian anyons have a quantum dimension of one. Even though we used the analogy to a tensor product of Hilbert spaces, note that the dimension of each Hilbert space is always an integer, while the quantum dimension is in general not. This is an important property of non-Abelian anyons that sets them apart from a simple set of particles inhabiting local Hilbert spaces.

²Note that in a system of $N > 2$ anyons we usually choose a basis of fully fused states that is obtained by successively fusing the two left-most anyons with respect to a fixed reference order $X_1 X_2 \dots X_N$. For $N = 3$ we would, for instance, write $|(X_1 X_2) X_3; (X_k X_3); X_l\rangle$, where k labels the all the possible fusion outcomes between X_1 and X_2 and where l labels all the possible fusion outcomes from $X_k X_3$ for every k , respectively.

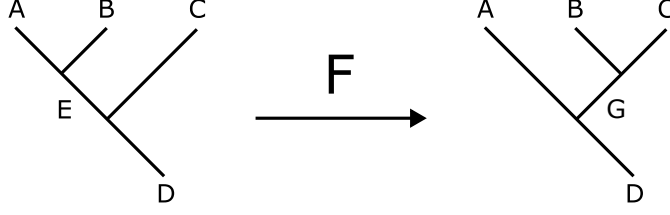


Figure 6.2: The two fusion diagrams associated with the fixed total charge D subspace of the fusion space.

The fusion space of ground states $\mathcal{H}_0$ is a collective non-local property of the non-Abelian anyons. Its non-locality means that no local perturbation can lift the degeneracy of $\mathcal{H}_0$, making it an ideal place to encode quantum information. It is important to note that the coherent superposition required for encoding a qubit is only possible within subspaces of the fusion space that belong to the same total topological charge sector, i.e. to subspaces spanned by fully fused states of the same total topological charge. For example, the fusion space of a pair X_1 and X_2 of non-Abelian anyons satisfying $X_1 \otimes X_2 = \mathbf{1} \oplus X_3$ cannot directly be used to encode a qubit because the two states $|(X_1 X_1); \mathbf{1}\rangle$ and $|(X_1 X_2); X_3\rangle$ belong to different total topological charge sectors $\mathbf{1}$ and X_3 . Encoding quantum information thus requires not only a two-dimensional fusion space, but a fusion space with at least two-dimensional *total charge sectors*. By extension, this means that we need more than two non-Abelian anyons for topological quantum computation. The basis of higher dimensional fusion spaces is given by the fully fused states obtained by successively fusing neighbouring anyons in a fixed reference order. A different fusion order is then equivalent to a change of basis. For every anyon model there exists a set of matrices that relate different bases. These so-called F -matrices result from the aforementioned modular data that defines the respective anyon model.³ In practice, they are obtained as the solutions to a set of consistency equations called the pentagon equations that we will not discuss in detail here. In order to illustrate how the F -matrices give structure to the fusion space, we consider a system of three anyons X_1, X_2, X_3 that are constrained always to fuse to X_4 . In the following, we relabel the anyons X_i as $X_1 = A$, $X_2 = B$, $X_3 = C$, and $X_4 = D$ to avoid clutter and improve readability. Fixing the reference order $X_1 X_2 X_3 = ABC$, there are two alternative fusion bases

$$|(AB)C; EC; D\rangle \quad \text{and} \quad |A(BC); AG; D\rangle, \quad (6.41)$$

corresponding to first-fusing (AB) into E followed by (EC) into D , and first-fusing BC into G followed by (AG) into D , respectively.⁴ The fusion diagram representation of these two bases is illustrated in Fig. 6.2. The unitary transformation F that implements the basis change between them acts as

$$|(AB)C; EC; D\rangle = \sum_G (F_{ABC}^D)_{EG} |A(BC); AG; D\rangle, \quad (6.42)$$

where $(F_{ABC}^D)_{EG}$ denote the matrix elements of the matrix F_{ABC}^D which describes the fixed total charge D fusion subspace of the system of the three anyons A, B and C . Given the F -matrices of an anyon model, the possible statistics are described by a collection of compatible interchange operators R . These can be obtained by solving another set of consistency equations known as hexagon equations [104]. In our example, a clockwise⁵ exchange of the first-fused anyons A and B corresponds to the unitary transformation

$$|(BA)C; EC; D\rangle = \sum_H (R_{AB})_{EH} |(AB)C; HC; D\rangle =: \sum_H R_{AB}^H \delta_{EH} |(AB)C; HC; D\rangle, \quad (6.43)$$

where H covers all the possible fusion outcomes of A and B and δ_{GE} is the Kronecker delta. The numbers R_{AB}^H describe the exchange of the first-fused anyons A and B in the fusion channel where they fuse to H next. The Kronecker delta δ_{EH} makes R_{AB} a diagonal unitary matrix. Accordingly, R merely assigns phase factors $R_{AB}^H = \exp(i\varphi_{AB}^H)$ depending only on the respective fusion channel H . This is a general

³The modular data of some known models can be found in Ref. [104].

⁴Of course, we assume that the fusion rules are compatible with these processes.

⁵If R is understood as a clockwise interchange of anyons, $R^\dagger$ corresponds to the counter-clockwise exchange.

theme in the R -matrices: when two anyons A and B are exchanged in a basis $| (AB)C; EC; D \rangle$ where they are fused first, the R -matrix acts as a diagonal matrix of fusion channel dependent phase factors. The exchange of two anyons in a basis where they are *not* fused first can be described using a combination of F and R matrices. Take, for example, the exchange of B and C in the basis from Eq. (6.41). We may implement this by first applying the $F_{ABC}^{D\dagger}$ matrix to change the basis from $| (AB)C; EC; D \rangle$ to $| A(BC); AG; D \rangle$, then acting with the diagonal matrix R_{BC} to exchange the now first-fused anyons B and C , and finally returning to the original basis using F_{ABC}^D . In formula, this process takes the form

$$\begin{aligned}
F_{ABC}^D R_{BC} F_{ABC}^{D\dagger} | (AB)C; EC; D \rangle &\stackrel{(\diamond)}{=} \sum_{X,Y,Z} (F_{ABC}^D)_{EX} (R_{BC})_{XY} (F_{ABC}^{D\dagger})_{YZ} | (AB)C; ZC; D \rangle \\
&= \sum_{X,Y} (F_{ABC}^D)_{EX} (R_{BC})_{XY} | A(BC); AY; D \rangle \\
&= \sum_{X,Y} (F_{ABC}^D)_{EX} R_{BC}^X \delta_{XY} | A(BC); AY; D \rangle \\
&= \sum_Y (F_{ABC}^D)_{EY} | A(CB); AY; D \rangle \\
&= | (AC)B; EB; D \rangle.
\end{aligned} \tag{6.44}$$

In the first line, we spelled out the F and $F^\dagger$ transformations. Then we inverted the F transformation given in Eq. (6.42) to carry out the action of $F^\dagger$ in the second line. Finally, we inserted in the explicit form of $(R_{BC})_{XY}$ from Eq. (6.43) in the third line and simplified the resulting expression as far as possible. All unitary braiding transformations that arise due to the pairwise interchange of anyons can be constructed in a similar fashion from the elementary F - and R -matrices.

6.3.1 Fibonacci Anyons

In order to describe a non-trivial anyonic model, we need at least one more anyon type besides the trivial anyon type $\mathbf{1}$. The Fibonacci anyonic system is one of the simplest possible anyonic models, as it has precisely that minimum number of two anyon types: the trivial type $\mathbf{1}$, and the non-trivial type τ . We get the (very short) list

$$\{\mathbf{1}, \tau\} \tag{6.45}$$

of anyon types. The anyons of type τ are called Fibonacci anyons – the reason for this name will soon become clear. Fibonacci anyons are self-dual, i.e. a particle of type τ has an antiparticle that is also of type τ . This tells us that for Fibonacci anyons, the distinction between anyon type and anyon representative is redundant. We will therefore write τ for both the Fibonacci anyon representative and type. The fusion rules of a Fibonacci anyonic system are

$$\mathbf{1} \otimes \mathbf{1} = \mathbf{1}, \quad \tau \otimes \mathbf{1} = \tau, \quad \tau \otimes \tau = \mathbf{1} \oplus \tau, \tag{6.46}$$

where the $\oplus$ in the final line reflects the two possible fusion channels between τ and itself. Say we have three (well-separated) τ anyons in a plane. We would like to know the possible fusion outcomes when all three anyons are brought together. When the first two τ 's are combined, we may get a type $\mathbf{1}$ or τ anyon. If the resulting anyon is of type $\mathbf{1}$, then the fusion with the third τ yields another anyon of type τ . If the resulting anyon is of type τ , then fusion with the third τ can produce an anyon of either type, $\mathbf{1}$ or τ . Hence the fusion result is not unique and even if we fix the total charge to τ , there are still two distinct fusion paths to get there. To analyse this, we write down the possible fusion paths of the three anyons in the form of fusion diagrams, in which the fusion happens one by one from left to right. Using the same notation as before, we get

$$|(\tau\tau)\tau; \tau\tau; \mathbf{1}\rangle, \quad |(\tau\tau)\tau; \mathbf{1}\tau; \tau\rangle, \quad |(\tau\tau)\tau; \tau\tau; \tau\rangle. \tag{6.47}$$

As anticipated, there are two total charges for this system: in the first we end up in $\mathbf{1}$ and in the second we end up in τ . We can see that the fusion sector with fixed total charge $\mathbf{1}$ is one-dimensional, while the

N	0	1	2	3	4	5	6
F_N^1	1	1	1	1	2	3	5
F_N^τ	0	1	1	2	3	5	8

Table 6.1: Sector dimensions F_N^1 and F_N^τ of fixed total charge **1** and τ for N Fibonacci anyons.

fusion sector with fixed total charge τ is two-dimensional. It is going to be instructive to see how these fixed total charge sectors evolve when we add another Fibonacci anyon to the system. A system of four Fibonacci anyons has the fusion paths

$$\begin{aligned} & |(\tau\tau)\tau\tau; (\mathbf{1}\tau)\tau; (\tau\tau); \mathbf{1}\rangle, \quad |(\tau\tau)\tau\tau; (\tau\tau)\tau; (\tau\tau); \mathbf{1}\rangle \\ & |(\tau\tau)\tau\tau; (\mathbf{1}\tau)\tau; (\tau\tau); \tau\rangle, \quad |(\tau\tau)\tau\tau; (\tau\tau)\tau; (\mathbf{1}\tau); \tau\rangle, \quad |(\tau\tau)\tau\tau; (\tau\tau)\tau; (\tau\tau); \tau\rangle, \end{aligned} \quad (6.48)$$

so the fixed total charge **1** sector is now two-dimensional, while the fixed total charge τ sector turns out three-dimensional. If we denote the different fixed total charge sector dimensions in a system of N Fibonacci anyons by F_N^1 and F_N^τ , we inductively find that $F_{N+1}^\tau = F_N^\tau + F_{N-1}^\tau$ for all $N \in \mathbb{N}$ with $N \geq 1$ and $F_{N+1}^1 = F_N^1 + F_{N-1}^1$ for all $N \in \mathbb{N}$ with $N \geq 3$. Table (6.1) lists the first few sector dimensions to illustrate the induction. After an appropriate index offset of $N = 1$ ($N = 3$) in the fixed total charge τ (**1**) sector, both sector dimensions evolve in accordance with the Fibonacci sequence. It is this feature to which Fibonacci anyons owe their name. Note that the Fibonacci evolution of sector dimensions suggests that the quantum dimension d_τ of the Fibonacci anyons is related to the golden ratio: the ground-state degeneracy F_N^τ of the fixed total charge τ sector in the presence of N Fibonacci anyons is given by the N -th element of the regular Fibonacci sequence, which is known to be asymptotic to $\phi^N/\sqrt{5}$, i.e. it fulfils

$$F_n^\tau \xrightarrow{N \rightarrow \infty} \frac{\phi^N}{\sqrt{5}} \propto \phi^N. \quad (6.49)$$

We can explicitly determine the quantum dimension d_τ of τ from the fusion rules in Eq. (6.46) as the largest⁶ eigenvalue of the τ fusion matrix $(M_\tau)_{jk} = M_{\tau j}^k$ with $j, k \in \{\tau, \mathbf{1}\}$. That fusion matrix reads

$$M_\tau = \begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix}, \quad (6.50)$$

so its eigenvalues are

$$\lambda_\pm = \frac{1}{2} \pm \sqrt{\frac{1+4}{4}} = \frac{(1 \pm \sqrt{5})}{2}. \quad (6.51)$$

The quantum dimension of τ is therefore

$$d_\tau := \frac{(1 + \sqrt{5})}{2}, \quad (6.52)$$

which is precisely the golden ratio ϕ . Importantly, the Fibonacci evolution of sector dimensions also points to a peculiarity of the Fibonacci fusion space: it lacks a natural tensor product structure in the sense that its dimensionality grows by an additive constant per τ anyon, rather than by a multiplicative one. As we will discuss shortly, this property of Fibonacci anyons causes some problems for their use in quantum computation. In order to encode a qubit in the fusion space of a system of Fibonacci anyons, we need to figure out a suitable fusion subspace. A qubit can be realised in any coherent two-dimensional subspace, which suggests we choose either the three anyon fusion subspace spanned by the fixed total charge τ states

$$|(\tau\tau)\tau; \mathbf{1}\tau; \tau\rangle \quad \text{and} \quad |(\tau\tau)\tau; \tau\tau; \tau\rangle \quad (6.53)$$

or the four anyon fusion subspace spanned by the fixed total charge **1** states

$$|(\tau\tau)\tau\tau; (\mathbf{1}\tau)\tau; (\tau\tau); \mathbf{1}\rangle \quad \text{and} \quad |(\tau\tau)\tau\tau; (\tau\tau)\tau; \tau\tau; \mathbf{1}\rangle. \quad (6.54)$$

⁶This eigenvalue is also called the Perron-Frobenius eigenvalue.

Both Eq. (6.53) and Eq. (6.54) span coherent two-dimensional Hilbert spaces, which may be used to accommodate a qubit. Within each subspace, the basis transformation matrix F and exchange matrix R of first-fused anyons take the form

$$F = \begin{pmatrix} \phi^{-1} & \phi^{-1/2} \\ \phi^{-1/2} & -\phi^{-1} \end{pmatrix} \quad \text{and} \quad R = \begin{pmatrix} e^{i4\pi/5} & 0 \\ 0 & e^{-i3\pi/5} \end{pmatrix}. \quad (6.55)$$

Here, ϕ is the again golden ratio. Finally, let us stress that the distinct simplicity of Fibonacci anyons is somewhat deceiving. Most importantly, and by far most amazingly, it turns out that arbitrary unitary transformations on the fusion subspaces can be implemented by braiding the Fibonacci anyons, which is why the Fibonacci anyon model is said to be *universal* for quantum computation. Looking at Eq. (6.55), this may come as a surprise: after all, we can easily tell that $R^{10} = \mathbb{1}$, i.e. that R is of finite order ten. Still, the braid group generated by R and $F^{-1}RF$ is dense in $SU(2)$, ensuring that the above statement is in fact true. Their universality makes Fibonacci anyons something like the Holy Grail of topological quantum computation.

Yet, Fibonacci quantum computation also has some limitations. Most importantly, the lack of a tensor product structure usually means that only *subspaces* of a given fixed total-charge sector are used to encode quantum information. For example, if three τ anyons are used to encode one qubit, then we would like to use twice that to encode two qubits. However, the fusion space for six τ anyons is eight-dimensional in the τ sector and five-dimensional in the $\mathbf{1}$ sector, so the logical qubit would only reside in a *subspace* of either sector. The second problem is that approximating even the simplest gates through repeated braiding is very complicated: even the NOT-gate requires thousands of braiding operations in a specific order [103]. Finally, Fibonacci anyons are difficult to find in nature. Until now, few microscopic systems were shown to support Fibonacci anyons at all and for the few systems that have been found in theory, it remains unclear whether they can ever be realised in a laboratory. For more details see Ref. [103].

6.3.2 Ising Anyons

The Fibonacci anyon model features two types of anyons, representing a minimal non-trivial anyon model. By contrast, the Ising anyon model adds complexity by considering a total three particle types: the trivial type $\mathbf{1}$, the non-trivial type ψ (fermion type) and the non-trivial type σ (Ising anyon type). This leaves us with the set

$$\{\mathbf{1}, \psi, \sigma\} \quad (6.56)$$

of anyon types. As a consequence, the fusion rules of an Ising anyon model are a bit more extensive. They read

$$\begin{aligned} \mathbf{1} \otimes \mathbf{1} &= \mathbf{1}, & \mathbf{1} \otimes \sigma &= \sigma, & \mathbf{1} \otimes \psi &= \psi \\ \psi \otimes \psi &= \mathbf{1}, & \psi \otimes \sigma &= \sigma, & \sigma \otimes \sigma &= \mathbf{1} \oplus \psi, \end{aligned} \quad (6.57)$$

where the $\oplus$ once more reflects the two possible fusion channels between σ and itself. The fusion rule $\psi \otimes \psi = \mathbf{1}$ indicates that, when brought together, two fermions behave the same as if there were no particle. Accordingly, $\psi \otimes \sigma = \sigma$ tells us that a ψ with a σ is indistinguishable from a single σ . The non-Abelian nature of the Ising anyons σ is encoded in the last fusion rule, which states that two Ising anyons can either behave as the trivial type $\mathbf{1}$ or as a fermion type ψ . Physically, these fusion rules can, for example, be realised in terms of topological p -wave superconductors. In that context, the trivial type $\mathbf{1}$ corresponds to the Bogoliubov vacuum, i.e. a condensate of Cooper pairs. The fermions ψ are complex Bogoliubov quasiparticles, which can combine into a Cooper pair and merge into the Bogoliubov vacuum. Finally, the Ising anyons σ emerge as defect-bound Majorana zero modes. As we will soon explain, there is a way in which a Majorana zero mode corresponds to *half* a complex Bogoliubov fermion. Since the actual physics happens in terms of these complex fermions, Majorana zero modes can only appear in even numbers in nature. A pair of defects can then carry a single non-local Bogoliubov fermion ψ that can either be unoccupied, which corresponds to the fusion channel $\sigma \otimes \sigma = \mathbf{1}$, or occupied, which corresponds to the fusion channel $\sigma \otimes \sigma = \psi$. We will come back to this in more detail later on. In a system of two

Ising anyons σ , we get fusion diagram basis states

$$|(\sigma\sigma); \mathbf{1}\rangle \quad \text{and} \quad |(\sigma\sigma); \psi\rangle, \quad (6.58)$$

so the fusion space for two Ising anyons is two-dimensional. However, either fixed total charge sector is only one-dimensional, so that, as before, we need at least three Ising anyons σ to encode a qubit. Since physical Ising anyons can only exist in pairs, we consider a system of four Ising anyons σ , finding

$$\begin{aligned} & |(\sigma\sigma)\sigma\sigma; (\mathbf{1}\sigma)\sigma; (\sigma\sigma); \mathbf{1}\rangle, \quad |(\sigma\sigma)\sigma\sigma; (\psi\sigma)\sigma; (\sigma\sigma); \mathbf{1}\rangle \\ & |(\sigma\sigma)\sigma\sigma; (\mathbf{1}\sigma)\sigma; (\sigma\sigma); \psi\rangle, \quad |(\sigma\sigma)\sigma\sigma; (\psi\sigma)\sigma; (\sigma\sigma); \psi\rangle, \end{aligned} \quad (6.59)$$

where the two fusion channels at the top correspond to fixed total charge $\mathbf{1}$, while the two fusion channels at the bottom correspond to fixed total charge ψ . We can surmise that the fusion space of Ising anyons has a natural tensor space structure: its dimension doubles for every added pair of Ising anyons σ . Hence, $2N$ Ising anyons σ yield a total fusion space of dimension 2^N . Equation (6.40) then allows us to immediately determine the Ising quantum dimension as

$$d_\sigma^{2N} \stackrel{!}{=} 2^N \implies d_\sigma \stackrel{!}{=} \sqrt{2}. \quad (6.60)$$

Indeed, explicit calculation shows that the eigenvalues of the Ising fusion matrix

$$M_\sigma = \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix}, \quad (6.61)$$

given in the basis $\{\mathbf{1}, \sigma, \psi\}$, are

$$\lambda_1 = 0, \quad \lambda_2 = -\sqrt{2}, \quad \lambda_3 = \sqrt{2}, \quad (6.62)$$

such that $d_\sigma = \max(0, -\sqrt{2}, \sqrt{2}) = \sqrt{2}$ confirms our initial guess. Now there is another non-trivial anyon type in the Ising anyon model, namely the fermion type ψ . If we calculate its quantum dimension d_ψ via the ψ fusion matrix

$$M_\psi = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix}, \quad (6.63)$$

and its eigenvalues

$$\lambda_1 = \lambda_2 = 1 \quad \text{and} \quad \lambda_3 = -1, \quad (6.64)$$

we find that the fermion type has quantum dimension $d_\psi = 1$ and is therefore Abelian. The only non-Abelian anyon in the Ising anyon model is the therefore Ising anyon itself. Of course, this information is already baked into the nomenclature of the model: the fermion type is called the fermion type because its anyonic statistics are doubly trivial in the sense that it is an Abelian mode with exchange phase $\theta = \pi$. We can encode a qubit in either one of the fusion subspaces with fixed total charge shown in Eq. (6.59). In both of these subspaces, the basis transformation matrix F and exchange matrix R of first-fused anyons take the form

$$F = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \quad \text{and} \quad R = e^{-i\pi/8} \begin{pmatrix} 1 & 0 \\ 0 & e^{i\pi/2} \end{pmatrix}. \quad (6.65)$$

Importantly, we find that a double exchange of the first-fused anyons becomes

$$R^2 = e^{-i\pi/4} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad (6.66)$$

which, up to the global phase factor of $\exp(-i\pi/4)$, represents the z -Pauli matrix σ_z . This allows us to identify the double exchange of first-fused Ising anyons as an implementation of a logical quantum Z -gate.

One can ask whether the double exchange of two anyons that do not form a first-fused pair results in something else. Following Eq. (6.44), we can examine this by calculating the the corresponding exchange matrix explicitly. For example, the double exchange of the second and third Ising anyons corresponds to the matrix

$$\begin{aligned} FR^2F^\dagger &= \frac{e^{-i\pi/4}}{2} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \\ &= \frac{e^{-i\pi/4}}{2} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ -1 & 1 \end{pmatrix} \\ &= e^{-i\pi/4} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \end{aligned} \tag{6.67}$$

which, up to the global phase factor of $\exp(-i\pi/4)$, represents the x -Pauli matrix σ_x , rather than the z -Pauli matrix we found before. The double exchange of anyons that do not form a first-fused pair therefore implements a logical quantum X -gate instead of a Z -gate. We conclude that the choice of the fusion basis, which corresponds to a choice of the Hilbert space basis in physical theories, determines the effect that elementary braiding processes between anyons have on the qubit subspace. Moreover, the braiding protocols for implementing the Z - and X -gates are remarkably simple, promising manageable experimental control requirements.

Alas, it turns out that the logical Z - and X -gates are the only logical gates, that can be implemented on a single Ising anyon qubit. This means that Ising anyons, in contrast to Fibonacci anyons, are far from universal for topological quantum computation. Despite this, Ising anyons are still the most promising candidates for testing topological quantum computation experimentally, as they are relatively easy to realise using defect bound Majorana zero modes in p -wave superconducting systems.

6.4 Emergent Anyons in Topological Superconductors

Consider a two-dimensional chiral $p_x + ip_y$ superconductor. A $\Phi = hcN/2e$ vortex is a topological defect in the superconducting order parameter $\Delta(\mathbf{r}) = \Delta_0(\mathbf{r})e^{i\phi(\mathbf{r})}$ [86, 105–107]. Around the vortex, the superconducting phase $\phi(\mathbf{r})$ winds by $2\pi N$, and this non-trivial winding enforces a phase singularity at the vortex core. To maintain a single-valued order parameter, the amplitude $\Delta_0(\mathbf{r})$ must vanish at the core, i.e. $\Delta(\mathbf{r}_{\text{core}}) = 0$. This local vanishing of the gap allows single quasiparticle states to appear in the core. The phase winding also induces a circulating supercurrent, which in turn traps the quantised flux $\Phi = hcN/2e$ and gives the Φ vortex its name. If the chiral $p_x + ip_y$ superconductor is either spinless, fully spin-polarised, or exhibits strong spin-orbit coupling, a quantum flux vortex can carry a self-adjoint Majorana zero mode (MZM) – a single quasiparticle state pinned to zero energy by particle-hole symmetry [86]. In the framework of the tenfold way of topological condensed matter [14], these MZMs emerge as topological defect modes, whose existence is guaranteed by the bulk topology of the ambient superconductor. Beyond their topological protection, MZMs also obey non-Abelian Ising anyon statistics. This property will be the focus of the following discussion.

For the next part we closely follow Refs. [95, 108, 109]. Take a two-dimensional chiral spinless $p_x + ip_y$ superconductor with two well-separated vortex-bound MZMs γ_A and γ_B . As before, the two MZM operators satisfy the Clifford algebra

$$\gamma_i = \gamma_i^\dagger \quad \text{and} \quad \{\gamma_i, \gamma_j\} = 2\delta_{ij}, \tag{6.68}$$

for $i, j \in \{A, B\}$ and combine into a single complex fermion

$$b_0 = \frac{1}{2}(\gamma_A + i\gamma_B) \tag{6.69}$$

that sits at zero energy. As a consequence, we get a two-dimensional subspace

$$\mathcal{H}_0 = \text{span}(|0\rangle, |1\rangle) \tag{6.70}$$

of degenerate many-body ground states spanned by

$$|0\rangle \quad \text{and} \quad |1\rangle \equiv b_0^\dagger |0\rangle. \quad (6.71)$$

Since $|0\rangle$ and $|1\rangle$ differ by one complex fermion, they have opposite fermion parities P_0 and P_1 . We will assume that the vortex-bound MZMs are exchanged when the vortices are interchanged. Thus, the exchange of vortex-bound MZMs amounts to a unitary transformation

$$U\gamma_A U^\dagger = \zeta_A \gamma_B \quad \text{and} \quad U\gamma_B U^\dagger = \zeta_B \gamma_A, \quad (6.72)$$

where the exchange phases ζ_i necessarily fulfill $\zeta_i \in \{-1, +1\}$ because of $\gamma_i = \gamma_i^\dagger$ and because the anticommutation relations must be preserved. However, the only unitary operators that exchange γ_A and γ_B while leaving everything else unchanged are of the form

$$\begin{aligned} U_\pm &= e^{\pm \frac{\pi}{4} \gamma_A \gamma_B} \\ &= \sum_n \frac{(\pm \frac{\pi}{4})^n}{n!} (\gamma_A \gamma_B)^n \\ &= \sum_n \frac{(\pm \frac{\pi}{4})^{2n}}{(2n)!} (\gamma_A \gamma_B)^{2n} + \sum_n \frac{(\pm \frac{\pi}{4})^{2n+1}}{(2n+1)!} (\gamma_A \gamma_B)^{2n+1} \\ &= \sum_n \frac{(\pm \frac{\pi}{4})^{2n}}{(2n)!} (-1)^{\lfloor \frac{2n}{2} \rfloor} + \sum_n \frac{(\pm \frac{\pi}{4})^{2n+1}}{(2n+1)!} (-1)^{\lfloor \frac{2n+1}{2} \rfloor} \gamma_A \gamma_B \\ &= \sum_n \frac{(\pm \frac{\pi}{4})^{2n}}{(2n)!} (-1)^n + \sum_n \frac{(\pm \frac{\pi}{4})^{2n+1}}{(2n+1)!} (-1)^n \gamma_A \gamma_B \\ &= \cos\left(\pm \frac{\pi}{4}\right) + \sin\left(\pm \frac{\pi}{4}\right) \gamma_A \gamma_B \\ &= \frac{1}{\sqrt{2}} (1 \pm \gamma_A \gamma_B), \end{aligned} \quad (6.73)$$

where we repeatedly used $\gamma_i^2 = 1$ and the anticommutation relations Eq. (6.68). The $\pm$ sign indicates the direction of the MZM exchange. Accordingly, the adjoint of $U_\pm$ is given by

$$U_\pm^\dagger = (e^{\pm \frac{\pi}{4} \gamma_A \gamma_B})^\dagger = e^{\pm \frac{\pi}{4} \gamma_B \gamma_A} = e^{\mp \frac{\pi}{4} \gamma_B \gamma_A} = U_\mp. \quad (6.74)$$

Under $U_\pm$ the MZM operators transform as

$$\begin{aligned} U_\pm \gamma_A U_\pm^\dagger &= e^{\pm \frac{\pi}{4} \gamma_A \gamma_B} \gamma_A e^{\mp \frac{\pi}{4} \gamma_A \gamma_B} \\ &= \left[\frac{1}{\sqrt{2}} (1 \pm \gamma_A \gamma_B) \right] \gamma_A \left[\frac{1}{\sqrt{2}} (1 \mp \gamma_A \gamma_B) \right] \\ &= \frac{1}{2} (1 \pm \gamma_A \gamma_B) (\gamma_A \mp \gamma_A^2 \gamma_B) \\ &= \frac{1}{2} (1 \pm \gamma_A \gamma_B) (\gamma_A \mp \gamma_B) \\ &= \frac{1}{2} (\gamma_A \mp \gamma_B \pm \gamma_A \gamma_B \gamma_A - \gamma_A \gamma_B^2) \\ &= \frac{1}{2} (\gamma_A \mp \gamma_B \mp \gamma_A^2 \gamma_B - \gamma_A) \\ &= \frac{1}{2} (\mp \gamma_B \mp \gamma_B) \\ &= \mp \gamma_B, \end{aligned} \quad (6.75)$$

and

$$\begin{aligned}
U_{\pm} \gamma_B U_{\pm}^{\dagger} &= e^{\pm \frac{\pi}{4} \gamma_A \gamma_B} \gamma_B e^{\mp \frac{\pi}{4} \gamma_A \gamma_B} \\
&= \left[\frac{1}{\sqrt{2}} (1 \pm \gamma_A \gamma_B) \right] \gamma_B \left[\frac{1}{\sqrt{2}} (1 \mp \gamma_A \gamma_B) \right] \\
&= \frac{1}{2} (1 \pm \gamma_A \gamma_B) (\gamma_B \mp \gamma_B \gamma_A \gamma_B) \\
&= \frac{1}{2} (1 \pm \gamma_A \gamma_B) (\gamma_B \pm \gamma_A \gamma_B^2) \\
&= \frac{1}{2} (\gamma_B \pm \gamma_A \pm \gamma_A \gamma_B^2 + \gamma_A \gamma_B \gamma_A) \\
&= \frac{1}{2} (\gamma_B \pm \gamma_A \pm \gamma_A - \gamma_A^2 \gamma_B) \\
&= \frac{1}{2} (\pm \gamma_A \pm \gamma_A) \\
&= \pm \gamma_A,
\end{aligned} \tag{6.76}$$

where we once again used $\gamma_i^2 = 1$ and the anticommutation relations Eq. (6.68). Combined, we get

$$\gamma_A \xrightarrow{U_{\pm}} \mp \gamma_B \quad \text{and} \quad \gamma_B \xrightarrow{U_{\pm}} \pm \gamma_A. \tag{6.77}$$

The important thing to notice is that γ_A and γ_B always acquire opposite signs under $U_{\pm}$, which is precisely the transformation behaviour of Ising anyons [101, 106]. To demonstrate that this is, in fact, the case, we use b_0 and $b_0^{\dagger}$ from Eq. (6.69) and write

$$\gamma_A = b_0^{\dagger} + b_0 \quad \text{and} \quad \gamma_B = i(b_0^{\dagger} - b_0). \tag{6.78}$$

If we plug this into Eq. (6.73) we get

$$\begin{aligned}
U_{\pm} &= e^{\pm \frac{\pi}{4} \gamma_A \gamma_B} \\
&= e^{\pm \frac{\pi}{4} (b_0^{\dagger} + b_0) i (b_0^{\dagger} - b_0)} \\
&= e^{\pm i \frac{\pi}{4} (b_0^{\dagger} b_0^{\dagger} - b_0^{\dagger} b_0 + b_0 b_0^{\dagger} - b_0 b_0)} \\
&= e^{\pm i \frac{\pi}{4} (b_0 b_0^{\dagger} - b_0^{\dagger} b_0)} \\
&= e^{\pm i \frac{\pi}{4} ([1 - b_0^{\dagger} b_0] - b_0^{\dagger} b_0)} \\
&= e^{\pm i \frac{\pi}{4}} e^{\mp i \frac{\pi}{2} b_0^{\dagger} b_0},
\end{aligned} \tag{6.79}$$

which we can write as

$$U_{\pm} = e^{\pm i \frac{\pi}{4}} e^{\mp i \frac{\pi}{2} n_0}, \tag{6.80}$$

where $n_0 = b_0^{\dagger} b_0$ is the number operator of the complex zero mode associated to b_0 . The two degenerate ground states $|0\rangle$ and $|1\rangle$ therefore transform as

$$U_{\pm} |0\rangle = e^{\pm i \frac{\pi}{4}} |0\rangle \quad \text{and} \quad U_{\pm} |1\rangle = e^{\pm i \frac{\pi}{4}} e^{\mp i \frac{\pi}{2}} |1\rangle. \tag{6.81}$$

Accordingly, the matrix representation of $U_{\pm}$ in the basis $\{|0\rangle, |1\rangle\}$ reads

$$\mathcal{U}_{\pm} = e^{\pm i \frac{\pi}{4}} \begin{pmatrix} 1 & 0 \\ 0 & e^{\mp i \frac{\pi}{2}} \end{pmatrix}. \tag{6.82}$$

Up to an overall phase factor, Eq. (6.82) is equal to the braiding matrix R of two first-fused algebraic Ising anyons given in Eq. (6.65). This suggests the identification

$$|0\rangle \simeq |(\sigma\sigma), \mathbf{1}\rangle, \quad |1\rangle \simeq |(\sigma\sigma), \psi\rangle \tag{6.83}$$

of the two physical ground states $|0\rangle$ and $|1\rangle$ of the topological superconductor with the fusion basis states $|(\sigma\sigma), \mathbf{1}\rangle$ and $|(\sigma\sigma), \psi\rangle$ of an algebraic two-Ising-anyon system, indicating an equivalence between

	γ_A	γ_B	b_0	$b_0^\dagger$	$ 0\rangle$	$ 1\rangle$	$\Delta\phi$
$U_\pm$	$\mp\gamma_B$	$\pm\gamma_A$	$\pm i b_0$	$\mp i b_0^\dagger$	$e^{\pm i \frac{\pi}{4}} 0\rangle$	$e^{\mp i \frac{\pi}{4}} 1\rangle$	$\mp \frac{\pi}{2}$
$U_\pm^2$	$-\gamma_A$	$-\gamma_B$	$-b_0$	$-b_0^\dagger$	$e^{\pm i \frac{\pi}{2}} 0\rangle$	$e^{\mp i \frac{\pi}{2}} 1\rangle$	$\mp \pi$
$U_\pm^3$	$\pm\gamma_B$	$\mp\gamma_A$	$\mp i b_0$	$\pm i b_0^\dagger$	$e^{\pm i \frac{3\pi}{4}} 0\rangle$	$e^{\mp i \frac{3\pi}{4}} 1\rangle$	$\mp \frac{3\pi}{2}$
$U_\pm^4$	γ_A	γ_B	b_0	$b_0^\dagger$	$e^{\pm i \pi} 0\rangle$	$e^{\mp i \pi} 1\rangle$	$\mp 2\pi$

Table 6.2: Transformation behaviour under multiple exchanges of γ_A and γ_B .

the vortex-bound MZMs and the algebraic Ising anyons themselves. However, the fact that Eqs. (6.82) and (6.65) differ by an overall phase factor seems to weaken this equivalence. Indeed, one can show that vortex-bound MZMs are only *projectively* equivalent to algebraic Ising anyons, which precisely means that the unitary braiding matrices U describing the anyon exchange in physical models are only equivalent to the algebraic braiding matrices up to a global, physically irrelevant phase factor [110–112]. The identification in Eq. (6.83) shows that the total charge **1** sector of the algebraic fusion space corresponds to the P_0 parity sector $\mathcal{H}_0^{(0)} = \text{span}(|0\rangle)$ of $\mathcal{H}_0$, while the total charge ψ sector corresponds to the P_1 parity sector $\mathcal{H}_0^{(1)} = \text{span}(|1\rangle)$.

Based on the unitary transformation induced by a single exchange of γ_A and γ_B , we can compute the outcomes for multiple exchanges. The results are listed in Tab. 6.2. Note that in this system the unitary exchange operator $U_\pm$ only causes an additional $U(1)$ phase factor for the physical states $|0\rangle$ and $|1\rangle$. Since $U(1)$ factors are Abelian, this mode of braiding is often called Abelian, even though the vortex-bound MZMs correspond to non-Abelian anyons. In order to fully access their non-Abelian statistics, we need a system with at least four MZMs $\gamma_A, \gamma_B, \gamma_C, \gamma_D$ corresponding to two complex zero energy fermions. Importantly, the way in which four MZMs combine into two complex fermionic states is no longer unique. In fact, this is true for every even number of $2N > 2$ MZMs, as discussed in greater detail in Sec. (5.5). For now, we choose a basis in which $\gamma_A, \gamma_B, \gamma_C, \gamma_D$ combine into complex fermion operators

$$b_{0,1} = \frac{1}{2}(\gamma_A + i\gamma_B) \quad \text{and} \quad b_{0,2} = \frac{1}{2}(\gamma_C + i\gamma_D). \quad (6.84)$$

Both $b_{0,1}$ and $b_{0,2}$ correspond to a complex fermion mode at zero energy, so we get a four-dimensional subspace

$$\mathcal{H}_0 = \text{span}(|0,0\rangle, |1,0\rangle, |0,1\rangle, |1,1\rangle), \quad (6.85)$$

spanned by degenerate many-body ground states

$$|n_1, n_2\rangle = b_{0,1}^{\dagger n_1} b_{0,2}^{\dagger n_2} |0\rangle, \quad (6.86)$$

which can be identified with fusion basis states

$$\begin{aligned} |0,0\rangle &\simeq |(\sigma\sigma)(\sigma\sigma), (\mathbf{11}), \mathbf{1}\rangle, & |1,1\rangle &\simeq |(\sigma\sigma)(\sigma\sigma), (\psi\psi), \mathbf{1}\rangle \\ |1,0\rangle &\simeq |(\sigma\sigma)(\sigma\sigma), (\psi\mathbf{1}), \psi\rangle, & |0,1\rangle &\simeq |(\sigma\sigma)(\sigma\sigma), (\mathbf{1}\psi), \psi\rangle. \end{aligned} \quad (6.87)$$

Note that the fusion order of the above fusion basis states is different from the one we used before: instead of fusing successively from left to right we now fuse from the outside in. This is tied to the basis choice we made in Eq. (6.84) because it prescribes $(\gamma_A \gamma_B)$ and $(\gamma_C \gamma_D)$ as the first-fused pairs of anyons. As before, the total charge **1** sector of the algebraic anyon model corresponds to the P_0 parity sector of the physical model, while the total charge ψ sector corresponds to the P_1 parity sector. Again, the exchange operator that implements an exchange of the MZMs γ_A and γ_B is given by

$$U_\pm^{(AB)} = e^{\pm i \frac{\pi}{4} \gamma_A \gamma_B} = e^{\pm i \frac{\pi}{4}} e^{\mp i \frac{\pi}{2} n_1}, \quad (6.88)$$

and it is easy to convince ourselves that it transforms

$$\gamma_A \xrightarrow{U_\pm^{(AB)}} \mp \gamma_B, \quad \gamma_B \xrightarrow{U_\pm^{(AB)}} \pm \gamma_A, \quad \gamma_C \xrightarrow{U_\pm^{(AB)}} \gamma_C, \quad \gamma_D \xrightarrow{U_\pm^{(AB)}} \gamma_D, \quad (6.89)$$

exchanging γ_A and γ_B , while leaving γ_C and γ_D unchanged. Of course, the exchange of γ_C and γ_D , which is implemented by

$$U_{\pm}^{(CD)} = e^{\pm i \frac{\pi}{4} \gamma_C \gamma_D} = e^{\pm i \frac{\pi}{4}} e^{\mp i \frac{\pi}{2} n_2}, \quad (6.90)$$

works analogously, transforming

$$\gamma_A \xrightarrow{U_{\pm}^{(CD)}} \gamma_A, \quad \gamma_B \xrightarrow{U_{\pm}^{(CD)}} \gamma_B, \quad \gamma_C \xrightarrow{U_{\pm}^{(CD)}} \mp \gamma_D, \quad \gamma_D \xrightarrow{U_{\pm}^{(CD)}} \pm \gamma_C. \quad (6.91)$$

The matrix representations of $U_{\pm}^{(AB)}$ and $U_{\pm}^{(CD)}$ on the subspace

$$\mathcal{H}_0 = \text{span}(|0, 0\rangle, |1, 0\rangle, |0, 1\rangle, |1, 1\rangle) \quad (6.92)$$

of degenerate eigenstates

$$|n_1, n_2\rangle = b_{0,1}^{\dagger n_1} b_{0,2}^{\dagger n_2} |0\rangle \quad (6.93)$$

are therefore

$$\mathcal{U}_{\pm}^{(AB)} = \begin{pmatrix} e^{\pm i \frac{\pi}{4}} & 0 & 0 & 0 \\ 0 & e^{\mp i \frac{\pi}{4}} & 0 & 0 \\ 0 & 0 & e^{\pm i \frac{\pi}{4}} & 0 \\ 0 & 0 & 0 & e^{\mp i \frac{\pi}{4}} \end{pmatrix} \quad \text{and} \quad \mathcal{U}_{\pm}^{(CD)} = \begin{pmatrix} e^{\pm i \frac{\pi}{4}} & 0 & 0 & 0 \\ 0 & e^{\mp i \frac{\pi}{4}} & 0 & 0 \\ 0 & 0 & e^{\mp i \frac{\pi}{4}} & 0 \\ 0 & 0 & 0 & e^{\pm i \frac{\pi}{4}} \end{pmatrix}, \quad (6.94)$$

which are given in the parity-sorted basis $\{|0, 0\rangle, |1, 1\rangle, |0, 1\rangle, |1, 0\rangle\}$. Recall that the requirement of coherent superpositions forces us to focus on fixed-parity sectors of $\mathcal{H}_0$ in topological quantum computation [102, 113]. If we restrict Eq. (6.94) to the fixed-parity subspaces $\mathcal{H}_0^{(0)} = \text{span}(|0, 0\rangle, |1, 1\rangle)$ with parity P_0 and $\mathcal{H}_0^{(1)} = \text{span}(|0, 1\rangle, |1, 0\rangle)$ with parity P_1 we get

$$\mathcal{U}_{\pm}^{(AB)}|_{P_0} = \begin{pmatrix} e^{\pm i \frac{\pi}{4}} & 0 \\ 0 & e^{\mp i \frac{\pi}{4}} \end{pmatrix} \quad \text{and} \quad \mathcal{U}_{\pm}^{(CD)}|_{P_0} = \begin{pmatrix} e^{\pm i \frac{\pi}{4}} & 0 \\ 0 & e^{\mp i \frac{\pi}{4}} \end{pmatrix} \quad (6.95)$$

and

$$\mathcal{U}_{\pm}^{(AB)}|_{P_1} = \begin{pmatrix} e^{\pm i \frac{\pi}{4}} & 0 \\ 0 & e^{\mp i \frac{\pi}{4}} \end{pmatrix} \quad \text{and} \quad \mathcal{U}_{\pm}^{(CD)}|_{P_1} = \begin{pmatrix} e^{\mp i \frac{\pi}{4}} & 0 \\ 0 & e^{\pm i \frac{\pi}{4}} \end{pmatrix}, \quad (6.96)$$

given in the bases $\{|0, 0\rangle, |1, 1\rangle\}$ and $\{|0, 1\rangle, |1, 0\rangle\}$, respectively. As expected, Eqs. (6.95) and (6.96) are again projectively equivalent to the double braiding matrix R of first-fused anyons, as given in Eq. (6.65). Accordingly, a double exchange of γ_A and γ_B induces a unitary transformation

$$\left(\mathcal{U}_{\pm}^{(AB)}|_{P_0}\right)^2 = \left(\mathcal{U}_{\pm}^{(AB)}|_{P_1}\right)^2 = e^{\pm i \frac{\pi}{2}} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} = e^{\pm i \frac{\pi}{2}} \sigma_z \quad (6.97)$$

on $\mathcal{H}_0^{(0)}$ and $\mathcal{H}_0^{(1)}$ that is projectively equivalent to R^2 from Eq. (6.66). The double exchange of γ_C and γ_D works analogously, the only difference being an inconsequential global phase factor of $e^{i\pi}$ on the P_1 parity sector:

$$\left(\mathcal{U}_{\pm}^{(CD)}|_{P_0}\right)^2 = e^{\pm i \frac{\pi}{2}} \sigma_z \quad \text{and} \quad \left(\mathcal{U}_{\pm}^{(CD)}|_{P_1}\right)^2 = e^{\mp i \frac{\pi}{2}} \sigma_z. \quad (6.98)$$

The fact that these braiding transformations are projectively equivalent to the z -Pauli matrix, tells us that a double exchange of γ_A and γ_B (γ_C and γ_D) can be used to implement a Z -gate on a qubit encoded in the fixed-parity subspaces. So far, we have only discussed the aforementioned Abelian mode of braiding, in which we exchange first-fused anyons and obtain diagonal transformations like the ones given in Eq. (6.94). In order to get *non-Abelian* braiding transformations, we have to consider exchanges of Ising anyons from different fusion pairs. Consider, for example, the exchange of γ_A and γ_D . As before, it is implemented by the unitary operator

$$U_{\pm}^{(AD)} = e^{\pm i \frac{\pi}{4} \gamma_A \gamma_D}. \quad (6.99)$$

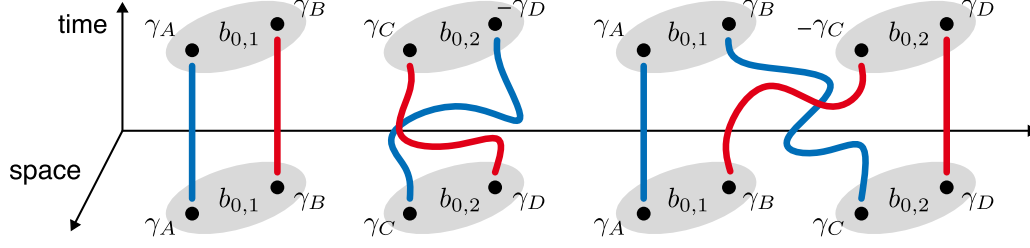


Figure 6.3: Two different modes of braiding MZMs. *Left:* “Abelian” braiding of two MZMs from the same complex fermion. *Right:* “non-Abelian” braiding of two MZMs from different complex fermions.

Using Eqs. (6.73) and (6.84), we can rewrite $U_{\pm}^{(AD)}$ as

$$U_{\pm}^{(AD)} = \frac{1}{\sqrt{2}}(1 \pm \gamma_A \gamma_D) = \frac{1}{\sqrt{2}}(1 \pm i(b_{0,1}^{\dagger} + b_{0,1})(b_{0,2}^{\dagger} - b_{0,2})), \quad (6.100)$$

which acts on a general $|n_1, n_2\rangle$ as

$$\begin{aligned} U_{\pm}^{(AD)} |n_1, n_2\rangle &= \frac{1}{\sqrt{2}}(1 \pm i(b_{0,1}^{\dagger} + b_{0,1})(b_{0,2}^{\dagger} - b_{0,2})) |n_1, n_2\rangle \\ &= \frac{1}{\sqrt{2}}(1 \pm i(b_{0,1}^{\dagger} b_{0,2}^{\dagger} - b_{0,1}^{\dagger} b_{0,2} + b_{0,1} b_{0,2}^{\dagger} - b_{0,1} b_{0,2})) |n_1, n_2\rangle \\ &= \frac{1}{\sqrt{2}} \begin{cases} (|1, 1\rangle \pm i |0, 0\rangle) & \text{for } |n_1, n_2\rangle = |1, 1\rangle \\ (|1, 0\rangle \mp i |0, 1\rangle) & \text{for } |n_1, n_2\rangle = |1, 0\rangle \\ (|0, 1\rangle \mp i |1, 0\rangle) & \text{for } |n_1, n_2\rangle = |0, 1\rangle \\ (|0, 0\rangle \pm i |1, 1\rangle) & \text{for } |n_1, n_2\rangle = |0, 0\rangle \end{cases} \\ &=: \frac{1}{\sqrt{2}}(|n_1, n_2\rangle \pm i(-1)^{\mathcal{N}} |1 - n_1, 1 - n_2\rangle), \end{aligned} \quad (6.101)$$

where we have used the standard action

$$\begin{aligned} c_i |n_1, \dots, n_i, \dots\rangle &= (-1)^{\sum_{j < i} n_j} \sqrt{n_i} |n_1, \dots, n_i - 1, \dots\rangle \\ c_i^{\dagger} |n_1, \dots, n_i, \dots\rangle &= (-1)^{\sum_{j < i} n_j} \sqrt{1 - n_i} |n_1, \dots, n_i + 1, \dots\rangle \end{aligned} \quad (6.102)$$

of fermionic annihilation and creation operators on Fock states and defined $\mathcal{N} := n_1 + n_2$ in the last line. The matrix representation of $U_{\pm}^{(AD)}$ is

$$\mathcal{U}_{\pm}^{(AD)} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \pm i & 0 & 0 \\ \pm i & 1 & 0 & 0 \\ 0 & 0 & 1 & \mp i \\ 0 & 0 & \mp i & 1 \end{pmatrix}, \quad (6.103)$$

which is given in the usual basis $\{|0, 0\rangle, |1, 1\rangle, |0, 1\rangle, |1, 0\rangle\}$. The restrictions of Eq. (6.103) to $\mathcal{H}_0^{(0)}$ with parity P_0 and $\mathcal{H}_0^{(1)}$ with parity P_1 take the form

$$\mathcal{U}_{\pm}^{(AD)}|_{P_0} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \pm i \\ \pm i & 1 \end{pmatrix} \quad \text{and} \quad \mathcal{U}_{\pm}^{(AD)}|_{P_1} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \mp i \\ \mp i & 1 \end{pmatrix}. \quad (6.104)$$

Consequently, a double exchange of γ_A with γ_D produces braiding transformations

$$\left(\mathcal{U}_{\pm}^{(AD)}|_{P_0}\right)^2 = \pm i \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \text{and} \quad \left(\mathcal{U}_{\pm}^{(AD)}|_{P_1}\right)^2 = \mp i \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad (6.105)$$

both of which are projectively equivalent to the x -Pauli matrix, such that a double exchange of γ_A and γ_D can be used to implement an X -gate on a qubit encoded in the fixed-parity subspaces. Note that there are four possible ways in which we can exchange MZMs from different complex fermions: we can

exchange γ_A and γ_C , γ_A and γ_D , γ_B and γ_C , and finally γ_B and γ_D . Simple symmetry considerations reveal, that the exchange of γ_C and γ_B is equivalent to the exchange of γ_A and γ_D . However, the exchange of γ_A and γ_C or that of γ_B and γ_D may produce different results. We test this explicitly, starting with the exchange of γ_A and γ_C . It is implemented by

$$U_{\pm}^{(AC)} = \frac{1}{\sqrt{2}}(1 \pm \gamma_A \gamma_C) = \frac{1}{\sqrt{2}}(1 \pm (b_{0,1} + b_{0,1}^\dagger)(b_{0,2} + b_{0,2}^\dagger)), \quad (6.106)$$

which results in a braiding transformation

$$U_{\pm}^{(AC)} |n_1, n_2\rangle = \frac{1}{\sqrt{2}}(|n_1, n_2\rangle \pm (-1)^{n_1} |1 - n_1, 1 - n_2\rangle) \quad (6.107)$$

with a matrix representation

$$\mathcal{U}_{\pm}^{(AC)} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \pm 1 & 0 & 0 \\ \mp 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & \pm 1 \\ 0 & \mp 1 & 1 & 1 \end{pmatrix}. \quad (6.108)$$

given in the same basis as before. The restrictions of $\mathcal{U}_{\pm}^{(AC)}$ to $\mathcal{H}_0^{(0)}$ and $\mathcal{H}_0^{(1)}$ are

$$\mathcal{U}_{\pm}^{(AC)}|_{P_0} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \pm 1 \\ \mp 1 & 1 \end{pmatrix} \quad \text{and} \quad \mathcal{U}_{\pm}^{(AC)}|_{P_1} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \pm 1 \\ \mp 1 & 1 \end{pmatrix}, \quad (6.109)$$

which square to

$$\left(\mathcal{U}_{\pm}^{(AC)}|_{P_0}\right)^2 = \mp i \sigma_y \quad \text{and} \quad \left(\mathcal{U}_{\pm}^{(AC)}|_{P_1}\right)^2 = \mp i \sigma_y. \quad (6.110)$$

The double exchange of these MZMs therefore implements a projective Y -gate. The exchange of γ_B and γ_D is analogous to that of γ_A and γ_C . The exchange operator in that case is

$$U_{\pm}^{(BD)} = \frac{1}{\sqrt{2}}(1 \mp (b_{0,1} - b_{0,1}^\dagger)(b_{0,2} - b_{0,2}^\dagger)), \quad (6.111)$$

which results in a braiding transformation

$$U_{\pm}^{(BD)} |n_1, n_2\rangle = \frac{1}{\sqrt{2}}(|n_1, n_2\rangle \mp (-1)^{n_2} |1 - n_1, 1 - n_2\rangle) \quad (6.112)$$

with a matrix representation

$$\mathcal{U}_{\pm}^{(BD)} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \mp 1 & 0 & 0 \\ \pm 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & \pm 1 \\ 0 & 0 & \mp 1 & 1 \end{pmatrix}, \quad (6.113)$$

The restrictions of $\mathcal{U}_{\pm}^{(BD)}$ to $\mathcal{H}_0^{(0)}$ and $\mathcal{H}_0^{(1)}$ are

$$\mathcal{U}_{\pm}^{(BD)}|_{P_0} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \mp 1 \\ \pm 1 & 1 \end{pmatrix} \quad \text{and} \quad \mathcal{U}_{\pm}^{(BD)}|_{P_1} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & \pm 1 \\ \mp 1 & 1 \end{pmatrix}, \quad (6.114)$$

which square to

$$\left(\mathcal{U}_{\pm}^{(BD)}|_{P_0}\right)^2 = \pm i \sigma_y \quad \text{and} \quad \left(\mathcal{U}_{\pm}^{(BD)}|_{P_1}\right)^2 = \mp i \sigma_y. \quad (6.115)$$

The double exchange of these MZMs therefore presents another way to implement a projective Y -gate.

6.5 Topological Quantum Computation with Anyons Revisited

In the limit of infinite anyon separation, anyon-anyon interactions become negligible and the fusion subspace $\mathcal{H}_0$ constitutes a truly degenerate subspace of the total Hilbert space. As a result, all states of $\mathcal{H}_0$ evolve according to the same unitary time evolution and there is no dynamical dephasing compromising its coherence. In combination with the topological protection against local perturbations and control errors, this inherent dynamical coherence makes topological quantum computation with anyons an exceptionally robust quantum computation scheme [114].

Of course, these ideal conditions are never quite met in the real world. In fact, the same microscopic mechanisms that give rise to anyonic quasiparticles in the first place also furnish them with interactions, introducing some degree of decoherence and perturbation-induced energy splitting to the system. Generally, these interactions decay exponentially in the separation between anyons, so it is important to keep anyonic quasiparticles “sufficiently far away” from one another in the laboratory. How far exactly “sufficiently far away” is, is determined by a coherence length l that is characteristic to the system. When two anyons A and B with two possible degenerate fusion channels $|(AB); X\rangle$ and $|(AB); Y\rangle$ are within a distance L of each other, the pairwise interaction between A and B lifts the degeneracy of $|(AB); X\rangle$ and $|(AB); Y\rangle$ by

$$\Delta\epsilon \sim e^{-L/l}, \quad (6.116)$$

where l denotes the parameter-dependent coherence length of the system. In practice, this lifting of the fusion channel degeneracy is always present to some extent. This leads to several subtleties for topological quantum computation. To begin with, the states of $\mathcal{H}_0$ will dephase over time, requiring a certain level of error correction despite the persistent topological protection [115]. The time after which substantial dephasing sets in is directly determined by the energy splitting $\Delta\epsilon$, highlighting once more the importance of anyon separation, even in finite-size limited environments. Next, a finite splitting of fusion channels means that the adiabatic exchange must not only be slow enough to avoid excitations beyond $\mathcal{H}_0$, but also fast enough to ensure that the fusion channels appear degenerate, see Sec. 4.5 and Refs. [116, 117]. Specifically, in finite-size systems there exist two characteristic time scales

$$\tau = \frac{\hbar}{\Delta\epsilon} \quad \text{and} \quad T = \frac{\hbar}{\Delta E} \quad (6.117)$$

associated to the largest energy splitting $\Delta\epsilon$ among the low-energy fusion channels and the energy gap ΔE between the highest energy state of $\mathcal{H}_0$ and the rest of the spectrum. Generally, we have $\Delta\epsilon \ll \Delta E$ such that $T \ll \tau$, i.e. the time scale τ that governs excitations within $\mathcal{H}_0$ is much larger (slower) than the time scale T that determines excitations from $\mathcal{H}_0$ to higher energy states. The time scale θ required for the adiabatic exchange of anyons with non-degenerate fusion channels must therefore satisfy

$$T \ll \theta \ll \tau, \quad (6.118)$$

i.e. it must be much faster than τ and much slower than T , cf. Sec. 4.5. Maintaining a balance between these two time scales generally leads to small errors in the implementation of quantum gates, which need to be corrected to avoid error accumulation and decoherence in the long run. Finally, interactions between anyons can even cause topological phase transitions, that destroy the topological protection and exotic statistical properties altogether [118–120]. The conditions under which such interaction-induced topological phase transitions occur have been the subject of several studies [121–124]. The severity and inevitability of these finite-size induced problems call for theoretical treatments and simulations that explicitly account for them. This will be one of the main topics in Chap. 10.

7 – Spectral Impurity Responses in Systems with Coexisting Topological Structures

It is one of the most remarkable insights in modern condensed matter physics, that the bulk topology of an insulator can enforce the existence of conducting states at its boundary. This profound connection is formalised in the celebrated tenfold way of topological quantum matter, which identifies entire symmetry classes of such materials [14]. The topology of these phases of matter is characterised by abstract invariants that capture the global properties of the underlying quantum states and may even manifest as experimentally accessible physical observables. The most famous example of the latter is the integer quantum Hall effect (IQHE), where the Hall conductivity, $\sigma_{xy} = -\frac{e^2}{h} C_1$, is determined by the first Chern number C_1 of the electronic Bloch bundle, allowing a direct experimental measurement of the system’s topological invariant and boundary modes [9, 10].

However, the immediate accessibility of topological invariants can vary greatly, so that alternative methods for the identification of topological states of matter become important. Apart from angle-resolved photoemission spectroscopy (ARPES), which can be used to image the surface band structure of three-dimensional topological materials [125], a particularly promising tool for the development of such methods could be local impurities. The possibility of probing the topology of a given quantum system using local impurities has been addressed in a number of studies [126–135]. These analyse the electronic structure near various types of impurities and impurity lattices in SPT systems across different Altland-Zirnbauer symmetry classes. The main goal of these works is to identify the general conditions under which the eigenenergies of the Hamiltonian robustly cross zero energy or traverse the band gap as external parameters are varied. In several lattice models, including the Kane–Mele [126], BHZ [131], and Haldane model [135], it was found that impurity-bound states appear in the topologically non-trivial phase but not in the trivial one.

In the following chapter, we study a two-dimensional lattice electron system with magnetic impurities at selected sites. To begin with, we consider a single magnetic impurity at an impurity site i_0 . It is modelled as a classical spin $\mathbf{S}$ and interacts with the local magnetic moment $\mathbf{s}_{i_0}$ of the electron system via a local exchange interaction $J\mathbf{s}_{i_0}\mathbf{S}$. For the electron model we choose the spinful Haldane model, which gives rise to a Bloch bundle over the Brillouin torus $\mathbb{T}_k^2$ that is characterised by the first Chern number $C_1^{(k)} \in \mathbb{Z}$, capturing the essential topological features of the integer quantum anomalous Hall effect (IQAHE). Additionally, the space of classical spin configurations forms a closed manifold $\mathcal{S} \equiv \mathbb{S}_S^2$ that enables a second, coexisting topological classification, but this time in terms of the first Chern number $Ch_1^{(S)} \in \mathbb{Z}$ over $\mathbb{S}_S^2$: if the Hamiltonian $H(\mathbf{S})$ depends smoothly on $\mathbf{S}$ and has a gapped, non-degenerate ground state for every $\mathbf{S} \in \mathbb{S}_S^2$, the associated $U(1)$ bundle of ground states admits a topological classification via the first Chern number [39, 136]. To distinguish between the Chern numbers of “ k -space” and “ S -space”, we denote them by $C_1^{(k)}$ and $Ch_1^{(S)}$, and call the latter the *spin*-Chern number. Just like in the conventional Bloch bundle topology, the spin-Chern number takes integer values and can change only when the gap closes, i.e. when ground-state degeneracies appear on a submanifold of $\mathbb{S}_S^2$.

The main goal of this work is to utilise the additional S -space topology to complement the conventional k -space classification. The interplay between both perspectives is expected to offer valuable insights into the behaviour of impurities within otherwise translationally invariant systems. This becomes particularly clear in the limit of strong exchange coupling, $J \rightarrow \infty$, where the local impurity physics is effectively governed by a magnetic-monopole model [2, 81, 137], $H_{\text{mono}} = J\mathbf{S}\mathbf{s}_{i_0}$, which exhibits a non-trivial S -space topology characterised by a non-trivial first spin-Chern number of $Ch_1^{(S)} = \pm 1$. Since the first spin-Chern number is trivially equal to zero for $J = 0$, it follows that there must be a spectral flow of impurity-bound in-gap states that crosses the chemical potential as a function of J . Moreover, we show that the inclusion of S -space topology gives rise to a non-trivial topological phase diagram in the critical interaction strength J_{crit} marking the S -space topological phase transition. This diagram reflects the coexisting k -space topology through its dependence on the parameters of the Haldane model.

Finally, we extend our study to the case of several impurity spins $\mathbf{S}_0, \dots, \mathbf{S}_{R-1}$, i.e. to a multi-impurity Kondo-Haldane model but with classical spins instead of localised quantum spins. With this modification, we focus on a regime where quantum-spin fluctuations and Kondo-screening effects can be disregarded. The configuration space of R classical impurity spins coupled to R distinct sites of the lattice is now the R -fold direct product $\mathcal{S}_R \equiv \mathbb{S}_{\mathbf{S}_0}^2 \times \dots \times \mathbb{S}_{\mathbf{S}_{R-1}}^2$ and we characterise the topology of ground state bundles over this $2R$ -dimensional manifold $\mathcal{S}_R$ by a characteristic number $Ch_R^{(S)}$, which we call the R -th spin-Chern number. As before, we trivially have $Ch_R^{(S)} = 0$ at $J = 0$, while we get $Ch_R^{(S)} = 1$ for $J \rightarrow \infty$. We find that the two distinct S -topological phases are separated by a finite range of coupling strengths, $J_{\text{crit},1} < J < J_{\text{crit},2}$, in which the system is gapless for at least one configuration $\mathbf{S}_0, \dots, \mathbf{S}_{R-1} \in \mathcal{S}_R$. The critical interactions $J_{\text{crit},1}$ and $J_{\text{crit},2}$ strongly depend on the parameters of the Haldane model and are found to be roughly one order of magnitude larger in the k -space topologically non-trivial phase than in the trivial phase. Systems with $R = 1$ and $R = 2$ impurity spins are studied numerically.

The remainder of this chapter is organised as follows. In Sec. 7.1 we introduce the concept of the spin-Chern number for electronic systems with classical spin impurities. Following this, we review the Haldane model and its k -space topology in Sec. 7.2. In the next section, Sec. 7.3, we present our results for the low-energy electronic structure in the presence of a single impurity spin. Section 7.4 provides an analysis of the spin-Chern number in the strong- J limit. This is followed by a discussion of the S -topological phase transition in Sec. 7.5. In Sec. 7.6 we examine how the S -space topological transition is affected by k -space topology. Finally, we evaluate numerical results for two impurity spins in the k -space trivial and non-trivial phases in Sec. 7.7.

Throughout this chapter, we closely follow our original presentation in [RQ1].

7.1 Multi-Impurity Kondo Model with Classical Spins

A multi-impurity Kondo model with R classical spins $\mathbf{S}_0, \dots, \mathbf{S}_{R-1}$ of fixed lengths $|\mathbf{S}_i| = 1$ for all $i = 0, \dots, R-1$ is described by a quantum-classical hybrid Hamiltonian

$$H(\mathbf{S}_0, \dots, \mathbf{S}_{R-1}) = H_{\text{el}} + H_{\text{int}}(\mathbf{S}_0, \dots, \mathbf{S}_{R-1}), \quad (7.1)$$

where H_{el} characterises any non-interacting lattice electron system, and

$$H_{\text{int}}(\mathbf{S}_0, \dots, \mathbf{S}_{R-1}) = J \sum_{m=0}^{R-1} \mathbf{S}_m \mathbf{s}_{i_m} \quad (7.2)$$

models the interactions between the $\mathbf{S}_m$ and the local magnetic moments $\mathbf{s}_{i_m}$ at the impurity sites i_m of the electronic lattice system via a local exchange coupling of strength J . Here, we choose an antiferromagnetic coupling $J > 0$. In second quantisation, H_{el} is constructed from the elementary fermionic creation and annihilation operators $c_{i\alpha}^\dagger$ and $c_{i\alpha}$ satisfying the canonical anticommutation relations

$$\{c_{i\alpha}, c_{j\beta}^\dagger\} = \delta_{ij} \delta_{\alpha\beta} \quad \text{and} \quad \{c_{i\alpha}, c_{j\beta}\} = \{c_{i\alpha}^\dagger, c_{j\beta}^\dagger\} = 0. \quad (7.3)$$

Here, i and j refer to sites of the underlying lattice and $\alpha, \beta \in \{\uparrow, \downarrow\}$ labels the electron spin projection. The components of the local magnetic moments $\mathbf{s}_{i_m}$ are given by

$$s_{i_m\mu} = \frac{1}{2} \sum_{\alpha, \beta} \sigma_\mu^{\alpha\beta} c_{i_m\alpha}^\dagger c_{i_m\beta}, \quad (7.4)$$

where $\mu = x, y, z$. The configuration space $\mathcal{S}_R$ of the R classical spins $\mathbf{S}_m \in \mathbb{S}_m^2$ is given by the R -fold direct product

$$\mathcal{S}_R \equiv \mathbb{S}_{\mathbf{S}_0}^2 \times \dots \times \mathbb{S}_{\mathbf{S}_{R-1}}^2 \quad (7.5)$$

and serves as an extrinsic parameter manifold for the model Eq. (7.1). Note that $\mathcal{S}_R$ is a simply-connected, closed, compact and orientable manifold of real dimension $\dim_{\mathbb{R}}(\mathcal{S}_R) = 2R$.

7.1.1 S -Space Topology of the Multi-Impurity Kondo Model with Classical Spins

The instantaneous eigenstates of Eq. (7.1) form complex vector bundles over $\mathcal{S}_R$ that may be analysed in terms of their characteristic classes. In particular, the many-body ground state $|\Psi_0(\mathbf{S}_0, \dots, \mathbf{S}_{R-1})\rangle$ of $H(\mathbf{S}_0, \dots, \mathbf{S}_{R-1})$ defines a complex line bundle $\Psi_0 \xrightarrow{\pi} \mathcal{S}_R$ whenever $|\Psi_0(\mathbf{S}_0, \dots, \mathbf{S}_{R-1})\rangle$ is non-degenerate and gapped on the entire parameter manifold $\mathcal{S}_R$.

We characterise this ground-state bundle based on a characteristic class $ch_R \in H^{2R}(\mathcal{S}_R, \mathbb{Q})$ which is known as the R -th Chern character of $\Psi_0 \xrightarrow{\pi} \mathcal{S}_R$, cf. Sec. 2.3.6. Recall that in Chern–Weil theory, ch_R is determined by the invariant polynomial

$$ch_R(\mathcal{F}) = \frac{1}{R!} \text{tr} \left(\frac{i\mathcal{F}}{2\pi} \right)^R \quad (7.6)$$

in the Berry curvature two-form

$$\mathcal{F} = \left(\frac{\partial \langle \Psi_0 |}{\partial S_\mu} \right) \left(\frac{\partial | \Psi_0 \rangle}{\partial S_\nu} \right) dS_\mu \wedge dS_\nu, \quad (7.7)$$

where μ and ν run over some choice of coordinates of the R impurity spins. Here, we wrote $|\Psi_0\rangle \equiv |\Psi_0(\mathbf{S}_0, \dots, \mathbf{S}_{R-1})\rangle$ for better readability. Since $ch_R(\mathcal{F}) \in H^{2R}(\mathcal{S}_R, \mathbb{Q})$ defines an element of the $2R$ -th rational cohomology group $H^{2R}(\mathcal{S}_R, \mathbb{Q})$ of $\mathcal{S}_R$, and $\mathcal{S}_R$ itself represents an element $\mathcal{S}_R \in H_{2R}(\mathcal{S}_R)$ of the $2R$ -th homology group $H_{2R}(\mathcal{S}_R)$, we can pair $ch_R(\mathcal{F})$ against $\mathcal{S}_R$ to obtain a characteristic rational number [39, 136]

$$Ch_R^{(S)} \equiv \langle [ch_R(\mathcal{F})], [\mathcal{S}_R] \rangle = \oint_{\mathcal{S}_R} ch_R(\mathcal{F}) = \frac{1}{R!} \left(\frac{i}{2\pi} \right)^R \oint_{\mathcal{S}_R} \text{tr}(\mathcal{F})^R, \quad (7.8)$$

called the R -th spin-Chern number $Ch_R^{(S)} \in \mathbb{Q}$ of $\Psi_0 \xrightarrow{\pi} \mathcal{S}_R$. Let us emphasise once more that the Chern character defines a cohomology class with rational coefficients. By this definition alone, $Ch_R^{(S)} \in \mathbb{Q}$ is therefore only guaranteed to be a rational number. The fact that Eq. (7.8), and more generally the top-degree Chern number $Ch_m = \langle [ch_m], [M] \rangle$ of any principal $U(1)$ -bundle over a $2m$ -dimensional basis manifold M , still yields an integer is a non-trivial result related to a deeper connection captured by the famous Atiyah–Singer index theorem. This was briefly addressed in Sec. 2.3.6.

We introduce coordinates $\boldsymbol{\lambda} = (\lambda_0, \dots, \lambda_{2R-1}) \equiv (\vartheta_0, \phi_0, \dots, \vartheta_{R-1}, \phi_{R-1})$ on $\mathcal{S}_R$, where each pair (ϑ_j, ϕ_j) denotes the polar and azimuthal angles on the j -th factor $\mathbb{S}_j^2$ of the product manifold. In these coordinates, the R -th spin-Chern number given in Eq. (7.8) can be computed as

$$Ch_R^{(S)} = \frac{1}{R!} \left(\frac{i}{2\pi} \right)^R \sum_{\pi \in S_{2R}} \text{sign}(\pi) \int d\lambda_0 \cdots d\lambda_{2R-1} \frac{\partial \langle \Psi_0 |}{\partial \lambda_{\pi(0)}} \frac{\partial | \Psi_0 \rangle}{\partial \lambda_{\pi(1)}} \cdots \frac{\partial \langle \Psi_0 |}{\partial \lambda_{\pi(2R-2)}} \frac{\partial | \Psi_0 \rangle}{\partial \lambda_{\pi(2R-1)}}, \quad (7.9)$$

where the sum runs over all permutations $\pi \in S_{2R}$ of the symmetric group of $2R$ elements. A few more details are presented in the appendix of [RQ1].

7.2 The Haldane Model

The Haldane model is an extended tight-binding model of graphene that captures the physics of the integer quantum anomalous Hall effect (IQAHE) [19]. It is based on the realisation that the only necessary condition for realising a topological QHE state is the breaking of TRS. Unlike conventional QHE systems, which rely on strong *external* magnetic fields to achieve this, Haldane designed his model to break TRS *intrinsically*. To achieve this, he introduced a periodic magnetic flux density $\mathbf{B}(\mathbf{r}) = B(\mathbf{r})\mathbf{e}_z$ with the full symmetry of the graphene lattice and zero flux through the unit cell. Before we move on to a discussion of the second-quantised Haldane Hamiltonian, we briefly go over the structural basics of the graphene lattice.

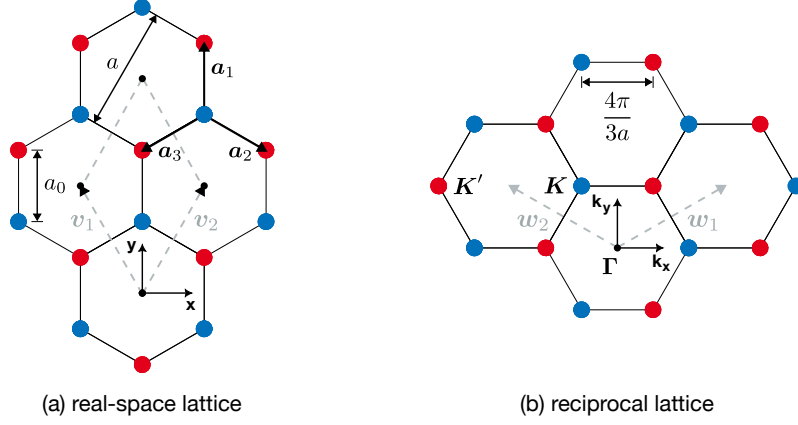


Figure 7.1: Sketches of the (a) real-space and (b) reciprocal honeycomb lattice. The blue and red dots label the sites of the two distinct sublattices A and B . The primitive lattice translation vectors $\mathbf{v}_1$ and $\mathbf{v}_2$ ($\mathbf{w}_1$ and $\mathbf{w}_2$) of the real-space (reciprocal) lattice are shown in grey colour. In the real-space lattice, the three nearest-neighbour positions $\mathbf{a}_1$, $\mathbf{a}_2$ and $\mathbf{a}_3$ are displayed and the interatomic distance a_0 and lattice constant a are indicated. In the reciprocal lattice, the positions the high-symmetry points Γ , K and K' are marked.

7.2.1 The Graphene Lattice

Graphene is a single layer of carbon atoms on a honeycomb lattice, as shown in Fig. 7.1. The honeycomb lattice Λ_h is not a Bravais lattice because its unit cell contains two carbon atoms, giving rise to two distinct sublattices A and B . For a honeycomb lattice in the xy -plane, the two primitive lattice translations are

$$\mathbf{v}_1 = \frac{a}{2} \begin{pmatrix} -1 \\ \sqrt{3} \end{pmatrix} \quad \text{and} \quad \mathbf{v}_2 = \frac{a}{2} \begin{pmatrix} 1 \\ \sqrt{3} \end{pmatrix}, \quad (7.10)$$

where $a = \sqrt{3}a_0 \approx 2.46 \text{ \AA}$ the lattice constant given in terms of the interatomic distance $a_0 \approx 1.42 \text{ \AA}$. The nearest-neighbour (NN) coordination number is three, while the next-nearest neighbour (NNN) coordination number is six. On sublattice A , the positions of the three NN sites are given by

$$\mathbf{a}_1 = a_0 \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \quad \mathbf{a}_2 = \frac{a_0}{2} \begin{pmatrix} \sqrt{3} \\ -1 \end{pmatrix}, \quad \mathbf{a}_3 = -\frac{a_0}{2} \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix}, \quad (7.11)$$

while those of the six NNN sites are take the form

$$\mathbf{b}_1 = -\mathbf{b}_4 = \mathbf{v}_1, \quad \mathbf{b}_2 = -\mathbf{b}_5 = \mathbf{v}_2, \quad \mathbf{b}_3 = -\mathbf{b}_6 = \mathbf{v}_2 - \mathbf{v}_1. \quad (7.12)$$

By definition, the honeycomb lattice is invariant under the action of the free Abelian group

$$T \simeq \text{FAB}(\mathbf{v}_1, \mathbf{v}_2) = \{ \mathbf{T} \in \mathbb{R}^2 \mid \mathbf{T} = n_1 \mathbf{v}_1 + n_2 \mathbf{v}_2, n_1, n_2 \in \mathbb{Z} \} \simeq \mathbb{Z}^2 \quad (7.13)$$

of lattice translations along $\mathbf{v}_1$ and $\mathbf{v}_2$. Its crystal structure is characterised by the point group D_{6h} [138]. For us, the most relevant subgroups of D_{6h} are a $C_6 \simeq \mathbb{Z}_6$ symmetry of rotations by $\phi_n = \frac{2\pi n}{6}$ ($n \in \mathbb{Z}$) around the center of each hexagon, a $I_2 \simeq \mathbb{Z}_2$ sublattice inversion symmetry of reflections about the center of each unit cell, and a $z \rightarrow -z$ reflection symmetry $R_2 \simeq \mathbb{Z}_2$ about the xy -plane.

In second quantisation, the periodic magnetic flux density of the Haldane model enters via a complex NNN Peierls hopping,

$$t_{\text{MF}} e^{\frac{ie}{\hbar} \int_{\mathbf{R}_k}^{\mathbf{R}_j} \mathbf{A}(\mathbf{r}) d\mathbf{r}} \equiv t_{\text{MF}} e^{i\xi_{jk}}, \quad (7.14)$$

where $\mathbf{R}_j$ and $\mathbf{R}_k$ denote the positions of the two NNN sites involved in the hopping process, $\mathbf{A}(\mathbf{r})$ is a suitable vector potential for $\mathbf{B}(\mathbf{r})$, and $\xi_{jk} \equiv \frac{e}{\hbar} \int_{\mathbf{R}_k}^{\mathbf{R}_j} \mathbf{A}(\mathbf{r}) d\mathbf{r}$ is the resulting Peierls phase. With this, the spinful Haldane Hamiltonian reads

$$H_H = -t_{\text{hop}} \sum_{\langle j,k \rangle} \sum_{\alpha} c_{j\alpha}^{\dagger} c_{k\alpha} + V \sum_{j,\alpha} \epsilon_j c_{j\alpha}^{\dagger} c_{j\alpha} + t_{\text{MF}} \sum_{\langle\langle j,k \rangle\rangle} e^{i\xi_{jk}} c_{j\alpha}^{\dagger} c_{k\alpha}, \quad (7.15)$$

$h_0(\mathbf{k})$	$2t_{\text{MF}} \cos \xi (2 \cos x \cos 3y + \cos 2x)$	$h_x(\mathbf{k})$	$-t_{\text{hop}} (2 \cos x \cos y + \cos 2y)$
$h_y(\mathbf{k})$	$-t_{\text{hop}} (2 \cos x \sin y - \sin 2y)$	$h_z(\mathbf{k})$	$V - 2t_{\text{MF}} \sin \xi (2 \sin x \cos 3y - \sin 2x)$

Table 7.1: Non-zero coefficients of Eq. (7.17), with $x = \sqrt{3}a_0k_x/2$ and $y = a_0k_y/2$. For details see App. A.8.

where j and k label the L sites of a honeycomb lattice Λ_h in the xy -plane, $\langle j, k \rangle$ and $\langle\langle j, k \rangle\rangle$ indicate summation over pairs of NN and NNN sites, and $\alpha, \beta \in \{\uparrow, \downarrow\}$ denote the spin projection of the electrons along the z -axis. The first term of Eq. (7.15) is the generic tight-binding hopping of graphene. It is governed by the real NN hopping amplitude t_{hop} (in graphene $t = 2.8\text{eV}$) and preserves all spatial (lattice and z -reflection) symmetries and the $\text{SU}(2)$ spin symmetry of the electrons. Moving forward, the lattice constant $a \equiv 1$ sets the length unit, while the NN hopping amplitude $t_{\text{hop}} \equiv 1$ sets the energy unit. The second term is a staggered sublattice potential. It is characterised by the real on-site potential strength V and the sign $\epsilon_j = \pm 1$, which is negative (positive) when j belongs to the A (B) sublattice. For $V \neq 0$, this term breaks the I_2 sublattice inversion symmetry ($I_2 \rightarrow 1$) and reduces the six-fold rotational symmetry C_6 to a three-fold rotational symmetry ($C_6 \rightarrow C_3$). The $\text{SU}(2)$ spin symmetry and the z -reflection symmetry R_2 are left invariant. The last term describes the coupling of the electrons to the periodic magnetic flux density. It is determined by the real NNN spin-orbit hopping amplitude t_{MF} and the real phase $\xi_{jk} = \pm \xi$, which is positive (negative) for anticlockwise (clockwise) hopping $k \rightarrow j$ within a hexagon of the lattice. Due to the complex phase factor, this term transforms as

$$\mathcal{T} t_{\text{MF}} \sum_{\langle\langle j, k \rangle\rangle_{\alpha}} e^{i\xi_{jk}} c_{j\alpha}^\dagger c_{k\alpha} \mathcal{T}^\dagger = -t_{\text{MF}} \sum_{\langle\langle j, k \rangle\rangle_{\alpha}} e^{-i\xi_{jk}} c_{j\alpha}^\dagger c_{k\alpha}, \quad (7.16)$$

under TRS. As a consequence, this term breaks TRS for $\xi \neq (2n+1)\pi/2$, allowing for a non-trivial QSH insulating state. A proof of Eq. (7.16) is given in App. A.8. We set $t_{\text{MF}} = 0.1$ throughout this chapter.

The Haldane Hamiltonian H_H can be diagonalised in $\mathbf{k}$ -space. Specifically, the Fourier transform $c_{j\alpha} = 1/\sqrt{L} \sum_{\mathbf{k}} e^{i\mathbf{k}\mathbf{R}_j} c_{\mathbf{k}\alpha}$ of the elementary field operators allows us to write H_H as

$$H_H = \sum_{\mathbf{k}} \phi(\mathbf{k})^\dagger h_H(\mathbf{k}) \phi(\mathbf{k}), \quad (7.17)$$

where we introduced the spinor $\phi(\mathbf{k}) = (a_{\mathbf{k}\uparrow} b_{\mathbf{k}\uparrow} a_{\mathbf{k}\downarrow} b_{\mathbf{k}\downarrow})^\top$ of annihilation operators $a_{\mathbf{k}\alpha}$ and $b_{\mathbf{k}\alpha}$ for Bloch states with quasi-momentum $\mathbf{k}$ and spin projection α on the A and B sublattices, respectively. Note that the A and B sublattices form a two-component degree of freedom that behaves mathematically like a spin. For this reason, it is often referred to as sublattice pseudospin. The spinors in Eq. (7.17) are then given in an $\sigma \otimes \tau$ tensor basis where σ is associated with the $\{\uparrow, \downarrow\}$ components of electron spin, while τ describes the $\{a, b\}$ components of sublattice pseudospin sector. The Hermitian 4×4 Bloch matrix $h_{\text{KM}}(\mathbf{k})$ from Eq. (8.3) takes the form (for details see App. A.8)

$$h_H(\mathbf{k}) = \mathbb{1}_2^\sigma \otimes [h_0(\mathbf{k}) \mathbb{1}_2^\tau + \mathbf{h}(\mathbf{k}) \boldsymbol{\tau}], \quad (7.18)$$

which is given in terms of the electron-spin (sublattice-pseudospin) identity matrix $\mathbb{1}_2^\sigma$ ($\mathbb{1}_2^\tau$) and the vector $\boldsymbol{\tau}$ of sublattice-pseudospin Pauli matrices. The coefficient functions $h_0(\mathbf{k})$ and $\mathbf{h}(\mathbf{k}) \equiv (h_x(\mathbf{k}) h_y(\mathbf{k}) h_z(\mathbf{k}))$ are listed in Tab. 7.1. The diagonalisation of Eq. (7.17) amounts to a diagonalisation

$$U(\mathbf{k})^\dagger h_H(\mathbf{k}) U(\mathbf{k}) = E_H(\mathbf{k}) = \text{diag}(E_\uparrow^+(\mathbf{k}), E_\downarrow^+(\mathbf{k}), E_\uparrow^-(\mathbf{k}), E_\downarrow^-(\mathbf{k})), \quad (7.19)$$

of the 4×4 Bloch matrix $h_H(\mathbf{k})$ from Eq. (7.18). The four energy bands

$$E_\alpha^\pm(\mathbf{k}) = h_0(\mathbf{k}) \pm \sqrt{h_x(\mathbf{k})^2 + h_y(\mathbf{k})^2 + h_z(\mathbf{k})^2} \equiv h_0(\mathbf{k}) \pm |\mathbf{h}(\mathbf{k})|, \quad (7.20)$$

exhibit a degeneracy between the spin-up ($\alpha = \uparrow$) and spin-down ($\alpha = \downarrow$) sectors. Thus, we get two doubly degenerate bands, a conduction band $E^+(\mathbf{k})$ and a valence band $E^-(\mathbf{k})$. These are symmetric around zero so the energy gap of the system is defined as

$$\Delta E := \min_{\mathbf{k}} (E^+(\mathbf{k}) - E^-(\mathbf{k})) = 2 \cdot \min_{\mathbf{k}} |\mathbf{h}(\mathbf{k})|. \quad (7.21)$$

We find that the minimum is attained at the Dirac points

$$\mathbf{K}^\pm = \pm \frac{4\pi}{3\sqrt{3}a_0} \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad (7.22)$$

where $h_x(\mathbf{K}^\pm) = h_y(\mathbf{K}^\pm) = 0$ and $h_z(\mathbf{K}^\pm) = V \mp 3\sqrt{3}t_{\text{MF}}\sin(\xi)$ so that

$$\Delta E := 2 \cdot \min_{\mathbf{K}^\pm} |V \mp 3\sqrt{3}t_{\text{MF}}\sin(\xi)|. \quad (7.23)$$

Note that $V \neq 0$ opens the gap symmetrically at both Dirac points, whereas $t_{\text{MF}} \neq 0$ modifies the gap size at $\mathbf{K}^\pm$ in opposite directions. Furthermore, the bulk gap ΔE closes along the nodal surface $|V| = 3\sqrt{3}|t_{\text{MF}}\sin(\xi)|$ in the three-dimensional parameter space spanned by t_{MF} , ξ and V . This nodal surface divides the parameter space into two separate types of regions with finite gaps $\Delta E > 0$: one dominated by the onsite potential ($|V| > 3\sqrt{3}|t_{\text{MF}}\sin(\xi)|$), and one dominated by the intrinsic spin-orbit coupling ($|V| < 3\sqrt{3}|t_{\text{MF}}\sin(\xi)|$). It turns out that these two regions are closely related to the topologically trivial and non-trivial phases of the Haldane model. For this reason, the onsite potential V is often used to tune between the topologically distinct phases in practice.

At half-filling, i.e. for every chemical potential μ with $|\mu| < \Delta E/2$, the ground state $|\text{GS}\rangle$ of the Haldane model is a Slater determinant

$$|\text{GS}\rangle = \prod_{\substack{\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^2 \\ \alpha = \uparrow, \downarrow}} d_{\mathbf{k}\alpha}^{-\dagger} |0\rangle = \bigwedge_{\substack{\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^2 \\ \alpha = \uparrow, \downarrow}} |u_{\alpha}^{-}(\mathbf{k})\rangle \quad (7.24)$$

of all valence Bloch states $|u_{\alpha}^{-}(\mathbf{k})\rangle \equiv d_{\mathbf{k}\alpha}^{-\dagger} |0\rangle \in \mathcal{H} \subset \mathcal{F}$. Here, $\mathcal{H} \subset \mathcal{F}$ indicates the natural inclusion of the single-particle Hilbert space $\mathcal{H}$ into the many-particle Fock space $\mathcal{F}$, and $|0\rangle$ denotes the vacuum state of $\mathcal{F}$ defined by $c_{\mathbf{k}\alpha} |0\rangle = 0$ for all $\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^2$ and $\alpha \in \{\uparrow, \downarrow\}$. The creation operators $d_{\mathbf{k}\alpha}^{-\dagger}$ of the valence Bloch states are defined along with creation operators $d_{\mathbf{k}\alpha}^{+\dagger}$ of the conduction Bloch states via

$$(d_{\mathbf{k}\uparrow}^{+\dagger}, d_{\mathbf{k}\downarrow}^{+\dagger}, d_{\mathbf{k}\uparrow}^{-\dagger}, d_{\mathbf{k}\downarrow}^{-\dagger}) := \phi(\mathbf{k})^{\dagger} U(\mathbf{k}), \quad (7.25)$$

i.e. via the diagonalisation transformation $U(\mathbf{k})$ of $h_{\text{H}}(\mathbf{k})$. The topology of the valence Bloch states therefore determines the topological properties of the many-body ground state.

7.2.2 Topology of the Haldane Model

The family $\{\mathcal{H}(\mathbf{k})\}_{\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^2}$ of Bloch spaces

$$\mathcal{H}(\mathbf{k}) := \text{span}(|u_{\uparrow}^{+}(\mathbf{k})\rangle, |u_{\downarrow}^{+}(\mathbf{k})\rangle, |u_{\uparrow}^{-}(\mathbf{k})\rangle, |u_{\downarrow}^{-}(\mathbf{k})\rangle) \quad (7.26)$$

defines a rank-four Bloch bundle $\mathcal{B}_{\text{H}} \xrightarrow{\pi} \mathbb{T}_{\mathbf{k}}^2$ over the two-dimensional Brillouin torus $\mathbb{T}_{\mathbf{k}}^2$. For $\Delta E > 0$, this Bloch bundle can be split as

$$\mathcal{B}_{\text{H}} = \mathcal{B}_{\text{H}}^{-} \oplus \mathcal{B}_{\text{H}}^{+}, \quad (7.27)$$

where $\mathcal{B}_{\text{H}}^{\pm} \xrightarrow{\pi^{\pm}} \mathbb{T}_{\mathbf{k}}^2$ are the two rank-two subbundles of $\mathcal{B}_{\text{H}}$ that are determined by the families $\{\mathcal{H}^{\pm}(\mathbf{k})\}_{\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^2}$ of valence and conduction subspaces

$$\mathcal{H}^{\pm}(\mathbf{k}) := \text{span}(|u_{\uparrow}^{\pm}(\mathbf{k})\rangle, |u_{\downarrow}^{\pm}(\mathbf{k})\rangle). \quad (7.28)$$

These are called the valence and conduction subbundles, accordingly. Since the spin-up and spin-down sectors of the spinful Haldane model in Eq. (7.15) are completely uncoupled, the rank-two valence and conduction subbundles can be further decomposed into a direct sum of rank-one line bundles

$$\mathcal{B}_{\text{H}}^{\pm} = \mathcal{B}_{\text{H},\uparrow}^{\pm} \oplus \mathcal{B}_{\text{H},\downarrow}^{\pm}, \quad (7.29)$$

so that the total Bloch bundle can be written as a direct sum

$$\mathcal{B}_{\text{H}} = \mathcal{B}_{\text{H},\uparrow}^{-} \oplus \mathcal{B}_{\text{H},\downarrow}^{-} \oplus \mathcal{B}_{\text{H},\uparrow}^{+} \oplus \mathcal{B}_{\text{H},\downarrow}^{+} \quad (7.30)$$

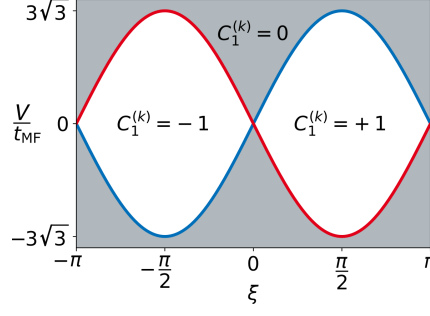


Figure 7.2: Phase diagram showing the first k -Chern number $C_1^{(k)} \equiv C_{1,(-,\uparrow)}^{(k)}$ of the spin-up valence bundle in the V/t_{MF} and ξ plane. Light grey and white colour highlight the trivial ($C_1^{(k)} = 0$) and non-trivial ($C_1^{(k)} = 1$) regions; red and blue lines mark gap closures at $\mathbf{K}^+$ and $\mathbf{K}^-$, see Ref. [19]. Reproduced with minor modifications from Ref. [RQ1].

of complex line bundles $\mathcal{B}_{\text{H},\alpha}^{\pm} \xrightarrow{\pi_{\alpha}^{\pm}} \mathbb{T}_{\mathbf{k}}^2$. We can characterise each of these rank-one subbundles using the Chern class $c_1 \in H^2(\mathbb{T}_{\mathbf{k}}^2, \mathbb{Z})$. In Chern–Weil theory, c_1 is determined by the invariant polynomial

$$c_1(\mathcal{F}_{\text{H},\alpha}^{\pm}) = \text{tr} \left(\frac{i\mathcal{F}_{\text{H},\alpha}^{\pm}}{2\pi} \right) \quad (7.31)$$

in the Berry curvature two-form

$$\mathcal{F}_{\text{H},\alpha}^{\pm} = \left(\frac{\partial \langle u_{\alpha}^{\pm}(\mathbf{k}) |}{\partial k_{\mu}} \right) \left(\frac{\partial |u_{\alpha}^{\pm}(\mathbf{k})\rangle}{\partial k_{\nu}} \right) dk_{\mu} \wedge dk_{\nu}, \quad (7.32)$$

where μ and ν denote any coordinates on $\mathbb{T}_{\mathbf{k}}^2$ and $|u(\mathbf{k})_{\alpha}^{\pm}\rangle$ is the Bloch state defining the valence $(-)$ or conduction $(+)$ subbundle with spin $\alpha \in \{\uparrow, \downarrow\}$. Since $c_1(\mathcal{F}_{\text{H},\alpha}^{\pm}) \in H^2(\mathbb{T}_{\mathbf{k}}^2, \mathbb{Z})$ defines an element of the second integral cohomology group $H^2(\mathbb{T}_{\mathbf{k}}^2, \mathbb{Z})$ of $\mathbb{T}_{\mathbf{k}}^2$, and $\mathbb{T}_{\mathbf{k}}^2$ itself represents an element $\mathbb{T}_{\mathbf{k}}^2 \in H_2(\mathbb{T}_{\mathbf{k}}^2)$ of the second homology group $H_2(\mathbb{T}_{\mathbf{k}}^2)$, we can pair $c_1(\mathcal{F}_{\text{H},\alpha}^{\pm})$ against $\mathbb{T}_{\mathbf{k}}^2$ to obtain a characteristic integer [39]

$$C_{1,(\pm,\alpha)}^{(k)} \equiv C_1^{(k)}(\mathcal{B}_{\text{H},\alpha}^{\pm}) = \langle [c_1(\mathcal{F}_{\text{H},\alpha}^{\pm})], [\mathbb{T}_{\mathbf{k}}^2] \rangle = \oint_{\mathbb{T}_{\mathbf{k}}^2} c_1(\mathcal{F}_{\text{H},\alpha}^{\pm}) = \frac{i}{2\pi} \oint_{\mathbb{T}_{\mathbf{k}}^2} \text{tr}(\mathcal{F}_{\text{H},\alpha}^{\pm}) \in \mathbb{Z}, \quad (7.33)$$

called the first k -Chern number $C_{1,(\pm,\alpha)}^{(k)}$ of $\mathcal{B}_{\text{H},\alpha}^{\pm} \xrightarrow{\pi_{\alpha}^{\pm}} \mathbb{T}_{\mathbf{k}}^2$. As we are dealing with Abelian $\text{U}(1)$ bundles, the trace in Eq. (7.33) is redundant and we arrive at the familiar expression

$$C_{1,(\pm,\alpha)}^{(k)} \equiv \frac{i}{2\pi} \oint_{\mathbb{T}_{\mathbf{k}}^2} \mathcal{F}_{\text{H},\alpha}^{\pm}. \quad (7.34)$$

The low dimensionality of the rank-one bundles makes it possible to compute these Chern numbers analytically. An exemplary calculation is included in App. A.8. Eventually, we get

$$C_{1,(-,\uparrow)}^{(k)} = C_{1,(-,\downarrow)}^{(k)} = -C_{1,(+,\uparrow)}^{(k)} = -C_{1,(+,\downarrow)}^{(k)} = \begin{cases} +1 & \text{for } |V| < 3\sqrt{3}|t_{\text{MF}} \sin(\xi)|, \sin(\xi) > 0 \\ 0 & \text{for } |V| > 3\sqrt{3}|t_{\text{MF}} \sin(\xi)| \\ -1 & \text{for } |V| < 3\sqrt{3}|t_{\text{MF}} \sin(\xi)|, \sin(\xi) < 0, \end{cases} \quad (7.35)$$

which is illustrated in Fig. 7.2 for $C_1^{(k)} \equiv C_{1,(-,\uparrow)}^{(k)}$. Equation (7.35) confirms that the total Bloch bundle $\mathcal{B}_{\text{H}}$ is always trivial, as

$$C_1^{(k)}(\mathcal{B}_{\text{H}}) = C_{1,(-,\uparrow)}^{(k)} + C_{1,(-,\downarrow)}^{(k)} + C_{1,(+,\uparrow)}^{(k)} + C_{1,(+,\downarrow)}^{(k)} = 0. \quad (7.36)$$

Meanwhile, the rank-two valence and conduction subbundles are non-trivial with

$$C_1^{(k)}(\mathcal{B}_{\text{H}}^{\pm}) = C_{1,(\pm,\uparrow)}^{(k)} + C_{1,(\pm,\downarrow)}^{(k)} = \pm 2 \quad (7.37)$$

in the topological phases. Finally, we use Eqs. (7.23) and (7.35) to define the critical sublattice potential strength

$$V_{\text{crit}} \equiv V_{\text{crit}}(t_{\text{MF}}, \xi) = 3\sqrt{3}|t_{\text{MF}} \sin(\xi)|, \quad (7.38)$$

at which the band gap closes and the k -Chern number becomes ill-defined.

7.2.3 Spin Haldane Model with Classical Spin Impurities

By adding the interaction term $H_{\text{int}}(\mathbf{S}_0, \dots, \mathbf{S}_{R-1})$ to the Haldane model, we break the translational symmetries. For $J \neq 0$, the k -space Chern number is therefore no longer well defined. Moreover, the classical spins act as local magnetic fields so that TRS is broken for $J \neq 0$ even if $t_{\text{MF}} = 0$. Note that the interaction term couples the two spin projections of the spinful Haldane Hamiltonian from Eq. (7.15). We also add a chemical-potential term $-\mu N$ to the Hamiltonian Eq. (7.15), where N is the total-particle number. Any value of the chemical potential μ inside the bulk band gap ensures a half-filled system. Its exact position within the gap, however, is relevant for the occupation of in-gap impurity states induced by the exchange interaction with the classical spins. As the impurity concentration R/L (here $R \leq 3$) is thermodynamically irrelevant, we set μ to its zero-temperature bulk value, i.e. μ lies exactly in the center of the bulk band gap. A different choice of μ will not qualitatively change the phase diagrams discussed later on. The total Hamiltonian from Eq. (7.1) can be cast into the form

$$H(\mathbf{S}_0, \dots, \mathbf{S}_{R-1}) = \sum_{j,k,\alpha,\beta} T_{(j,\alpha)(k,\beta)}^{\text{eff}}(\mathbf{S}_0, \dots, \mathbf{S}_{R-1}) c_{j\alpha}^\dagger c_{k\beta}, \quad (7.39)$$

where

$$T_{(j,\alpha)(k,\beta)}^{\text{eff}}(\mathbf{S}_0, \dots, \mathbf{S}_{R-1}) = T_{(j\alpha)(k\beta)} + \frac{J}{2} \sum_{m=0}^{R-1} (\mathbf{S}_m \boldsymbol{\sigma})^{\alpha\beta} \delta_{ji_m} \delta_{ki_m} \quad (7.40)$$

are the elements of the effective hopping matrix $T^{\text{eff}}(\mathbf{S}_0, \dots, \mathbf{S}_{R-1})$, given in terms of the elements

$$T_{(j\alpha)(k\beta)} = -t_{\text{hop}} \delta_{\langle j,k \rangle} \delta_{\alpha\beta} + V \epsilon_j \delta_{jk} \delta_{\alpha\beta} + t_{\text{MF}} e^{i\xi_{jk}} \delta_{\langle\langle j,k \rangle\rangle} \delta_{\alpha\beta} \quad (7.41)$$

of the hopping matrix of the pristine Haldane T model and the vector $\boldsymbol{\sigma}$ of electron-spin Pauli matrices. The single-particle energies $\varepsilon_n(\mathbf{S}_0, \dots, \mathbf{S}_{R-1})$ are obtained by numerical diagonalisation of $T^{\text{eff}}(\mathbf{S}_0, \dots, \mathbf{S}_{R-1})$ for arbitrary spin configurations.

7.3 Low-Energy Electronic Structure with a Single Impurity Spin

As a start, we consider a model with a single impurity spin, introducing the notation $\mathbf{S} \equiv \mathbf{S}_0$ for better readability. We determine the energy spectrum via exact diagonalisation of the full hopping matrix $T^{\text{eff}}(\mathbf{S})$, Eq. (7.40), and analyse it as a function of the local exchange-coupling strength J . Due to periodic boundary conditions, the results are independent of the choice of the unit cell. However, they still differ depending on whether the impurity spins are coupled to sites of sublattice A or B. The calculations presented here were done for a system, where the impurity spin is coupled to a sublattice A site i_0 of the honeycomb lattice.

Figure 7.3 shows the single-particle energies ε_n for generic parameters $\xi = \pi/4$ and $t_{\text{MF}} = 0.1$ within a narrow window of width $W_0 = 1.2$ around $\mu \approx -0.21$. For comparison, the valence and conduction bands are much broader, with widths $W_{\text{val}} \approx 2.2$ and $W_{\text{cond}} \approx 3.5$, respectively. The total width of the electronic band structure, including the band gap of about $\Delta E \approx 0.37$, is given by $W \approx 6.1$. Aside from the bulk band gap,

$$\Delta E = 2|V - 3\sqrt{3}t_{\text{MF}} \sin(\xi)| = 2|V - V_{\text{crit}}|, \quad (7.42)$$

and numerous small gaps caused by the finite lattice size $L = 2 \cdot 39^2 = 3042$, Fig. 7.3 highlights the presence of in-gap states and their dependence on the exchange-coupling strength J . We have considered three different sublattice potential strengths V to illustrate that the energies of the in-gap states depend strongly on the model of the underlying electron system. To facilitate direct comparison, the three sublattice potential strengths were chosen such that the bulk band gap ΔE is identical in all cases. Note that the ground states for $V = -0.5V_{\text{crit}}$ and $V = +0.5V_{\text{crit}}$ differ in the occupation of the impurity site,

$$n_{i_0} = \sum_{\alpha=\uparrow,\downarrow} c_{i_0\alpha}^\dagger c_{i_0\alpha}, \quad (7.43)$$

giving $\langle n_{i_0} \rangle > 1$ and $\langle n_{i_0} \rangle < 1$, respectively. Importantly, the systems with $V = \pm 0.5V_{\text{crit}}$, shown in the left and middle panel, are in a k -space topologically non-trivial state with $C_1^{(k)} = +1$, whereas the system with $V = 1.5V_{\text{crit}}$, shown in the right panel, is in the k -space topologically trivial state with $C_1^{(k)} = 0$.

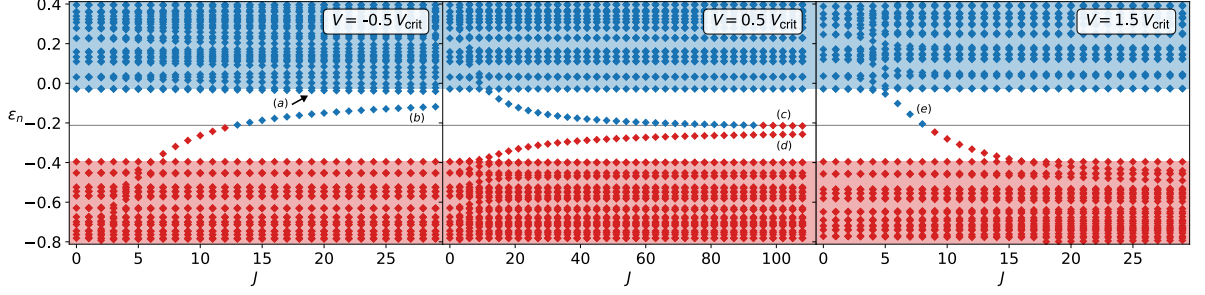


Figure 7.3: Single-particle energies as a function of J , as obtained by diagonalisation of the effective hopping matrix from Eq. (7.40) for a periodic honeycomb lattice with of 39×39 unit cells. A single spin $\mathbf{S}$ is coupled to sublattice A at a site i_0 . Calculations for sublattice potentials $V = \pm 0.5 V_{\text{crit}}$ (k -space topologically non-trivial) and $V = 1.5 V_{\text{crit}}$ (k -space topologically trivial), where $V_{\text{crit}} = 3\sqrt{3}t_{\text{MF}} \sin \xi$. Further parameters are $t_{\text{hop}} = 1$, $\xi = \pi/4$, $t_{\text{MF}} = 0.1$. The chemical potential $\mu \approx -0.21$ is located in the middle of the bulk band gap (gray line). Only the low-energy electronic structure is displayed with occupied (red) and unoccupied states (blue). In-gap bound states are labelled by (a) – (e). For each of the three considered sublattice potentials, an S -space topological transition occurs at a critical interaction strength J_{crit} , where an in-gap state, namely (b), (c), or (e), crosses the chemical potential, indicated by a colour change from red to blue or vice versa. Note that the middle panel displays a larger J -range to illustrate the convergence of the Zeeman pair inside the gap. Adapted with minor modifications from Ref. [RQ1].

Additionally, the $R = 1$ spin-Chern number was computed numerically from Eqs. (7.8) and (7.9). For $J = 0$, it vanishes in all configurations, $Ch_1^{(S)} = 0$, as the classical spin manifold $\mathcal{S}_1 = \mathbb{S}^2$ is entirely decoupled from the electron system. Increasing J eventually induces a topological transition, marked by a sudden jump to $Ch_1^{(S)} = 1$ at a critical exchange coupling J_{crit} . The system remains in this phase until the strong-coupling limit, $J \rightarrow \infty$, is realised. At $J = J_{\text{crit}}$, an in-gap single-particle state crosses the chemical potential, causing the many-body energy gap to close and the many-body ground state to become degenerate. Since $H_{\text{el}} + H_{\text{int}}(\mathbf{S})$ lacks two-electron interaction terms, the two degenerate many-electron ground states $|\Psi_{0,1}\rangle$ and $|\Psi_{0,2}\rangle$ are Slater determinants,

$$|\Psi_{0,1}\rangle = \prod_{\varepsilon_n < \mu} d_n^\dagger |0\rangle = \bigwedge_{\varepsilon_n < \mu} |\varepsilon_n\rangle \quad \text{and} \quad |\Psi_{0,2}\rangle = \prod_{\varepsilon_n \leq \mu} d_n^\dagger |0\rangle = \left[\bigwedge_{\varepsilon_n < \mu} |\varepsilon_n\rangle \right] \wedge |\mu\rangle, \quad (7.44)$$

differing in the occupation of the in-gap single-particle state $|\mu\rangle$. Accordingly, a different choice of μ would also shift the value of J_{crit} .

The quantum-classical Hamiltonian in Eq. (7.1) is $\text{SO}(3)$ symmetric. It is invariant under simultaneous rotations of the classical spins and the quantum spin degrees of freedom about an arbitrary axis $\mathbf{n}$ by an angle φ . In the classical sector, an $\text{SO}(3)$ rotation is represented by a matrix

$$O_{\mathbf{n}}(\varphi) = \exp(\mathbf{T}\mathbf{n}\varphi), \quad (7.45)$$

acting in the spin-configuration space $\mathcal{S}_1 = \mathbb{S}^2$. Here, $\mathbf{T} = (T_x \ T_y \ T_z)$ is the vector of real and skew-symmetric 3×3 matrices T_i generating the Lie algebra $\mathfrak{o}(3)$ of $\text{SO}(3)$. Consequently, the T_i satisfy

$$[T_i, T_j] = \epsilon_{ijk} T_k, \quad (7.46)$$

where we used an Einstein notation (implicit sum over k) for better readability. A common choice of generators is

$$T_x = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{pmatrix}, \quad T_y = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ -1 & 0 & 0 \end{pmatrix}, \quad T_z = \begin{pmatrix} 0 & -1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}. \quad (7.47)$$

In the quantum sector, the rotation is represented by a unitary operator $U_{\mathbf{n}}(\varphi) = \exp(-i\mathbf{s}_{\text{tot}}\mathbf{n}\varphi)$ with the total electron spin $\mathbf{s}_{\text{tot}} = \sum_i \mathbf{s}_i$. An immediate consequence of the invariance

$$U_{\mathbf{n}}(\varphi)H(O_{\mathbf{n}}(\varphi)\mathbf{S})U_{\mathbf{n}}^\dagger(\varphi) = H(\mathbf{S}) \quad (7.48)$$

is a continuous degeneracy of the eigenenergies. Specifically, the single-particle energies, and by extension the N -particle Fock state energies, are independent of the explicit orientation of the classical impurity spin $\mathbf{S}$. In particular, the single-particle energy of an in-gap state actually represents the energy of all single-particle states on the entire configuration manifold $\mathcal{S}_1 = \mathbb{S}^2$. As a result, one can assign a *single-particle* spin-Chern number $\chi_1^{(S)}$ to each individual in-gap state. The total spin-Chern number $C_1^{(S)}$ of the ground state corresponds to the sum of the single-particle spin-Chern numbers $\chi_1^{(S)}$ of the occupied states. Consequently, $C_1^{(S)}$ changes at J_{crit} when the in-gap state crossing the chemical potential carries a finite $\chi_1^{(S)} \neq 0$. As J increases, the total spin-Chern number changes by $\Delta Ch_1^{(S)} = -\chi_1^{(S)}$ if the in-gap state crosses the chemical potential from below (becomes unoccupied), and by $\Delta Ch_1^{(S)} = +\chi_1^{(S)}$ if it crosses from above (becomes occupied). With this, we find $\chi_1^{(S)} = +1$ for state (a), $\chi_1^{(S)} = -1$ for state (b), $\chi_1^{(S)} = +1$ for state (c), $\chi_1^{(S)} = -1$ for state (d), and $\chi_1^{(S)} = +1$ for state (e) in Fig. 7.3. All these configurations result in a net change of $\Delta Ch_1^{(S)} = +1$ in the total spin-Chern number. Note that the change $\Delta Ch_1^{(S)}$ is independent of the explicit choice of μ . For instance, lowering μ would cause state (d) with $\chi_1^{(S)} = -1$ (middle panel) to cross from below, while state (c) with $\chi_1^{(S)} = +1$ would remain unoccupied, instead of crossing from above. Accordingly, this configuration would result in the same net change of $\Delta Ch_1^{(S)} = +1$.

7.4 Strong- J Limit: The Magnetic Monopole

In the strong- J limit, the electronic system forms two impurity-bound high-energy states that can be viewed as the spin-down and spin-up states, $|\varepsilon_{i_0\downarrow}\rangle$ and $|\varepsilon_{i_0\uparrow}\rangle$, of a local Zeeman pair.¹ As $J \rightarrow \infty$, the energy of $|\varepsilon_{i_0\downarrow}\rangle$ approaches negative infinity, $\varepsilon_{i_0\downarrow} \rightarrow -\infty$, and the state becomes fully occupied. At the same time, the energy of $|\varepsilon_{i_0\uparrow}\rangle$ grows indefinitely, $\varepsilon_{i_0\uparrow} \rightarrow +\infty$, and the state is effectively removed from the system. As a result, the local occupation, $n_{i_0} = n_{i_0\uparrow} + n_{i_0\downarrow}$, tends towards half-filling, $n_{i_0} \rightarrow 1$, and the local magnetic moment $\langle \mathbf{s}_{i_0}^2 \rangle$ increasingly resembles that of a rigid quantum spin-1/2 with $\langle \mathbf{s}_{i_0}^2 \rangle \rightarrow 3/4$. The local physics for $J \rightarrow \infty$ is therefore effectively governed by the Hamiltonian

$$H_{\text{mono}}(\mathbf{S}) = J\mathbf{S}\mathbf{s}_{i_0}, \quad (7.49)$$

describing a rigid spin-1/2 $\mathbf{s}_{i_0}$ in an external magnetic field $J\mathbf{S}$. Remarkably, this system provides a paradigmatic realisation of a magnetic monopole [2, 81, 137]. To see this, we plug in the definition

$$\mathbf{s}_{i_0} := \frac{1}{2}\boldsymbol{\sigma}_{i_0} \quad (7.50)$$

of the quantum spin-1/2 operator $\mathbf{s}_{i_0}$ in terms of the Pauli vector $\boldsymbol{\sigma}_{i_0}$ and write

$$H(\mathbf{B}) = \mathbf{B}\boldsymbol{\sigma}_{i_0} = \begin{pmatrix} B_z & B_x - iB_y \\ B_x + iB_y & -B_z \end{pmatrix}, \quad (7.51)$$

where we defined the rescaled magnetic field

$$\mathbf{B} := \frac{J}{2}\mathbf{S} \quad (7.52)$$

to tidy up the notation. Equation (7.51) constitutes a generic two-level system with eigenvalues

$$E_{\pm} = \pm|\mathbf{B}| \equiv \pm B, \quad (7.53)$$

¹Note that, in this case, spin-up ($\uparrow$) and spin-down ($\downarrow$) refer to the local quantisation axis determined by the orientation of the classical impurity spin.

and normalised eigenstates (for details see App. A.9)

$$|\psi_{\pm}(\mathbf{B})\rangle = \frac{1}{\sqrt{2(\pm B)(B_z \pm B)}} \begin{pmatrix} B_z \pm B \\ B_x + iB_y \end{pmatrix}. \quad (7.54)$$

For $B \neq 0$, these states define two independent principal $U(1)$ bundles $\psi_{\pm} \xrightarrow{\pi_{\pm}} \mathbb{R}^3 \setminus \{\mathbf{0}\}$. As an example, consider the bundle $\psi_- \xrightarrow{\pi_-} \mathbb{R}^3 \setminus \{\mathbf{0}\}$ associated with the negative energy eigenstate. Straightforward but tedious calculation yields the Berry connection

$$\mathcal{A}_-(\mathbf{B}) = \langle \psi_-(\mathbf{B}) | d | \psi_-(\mathbf{B}) \rangle = \frac{i}{2} \frac{B_y dB_x - B_x dB_y}{B(B_z - B)}. \quad (7.55)$$

and the Berry curvature

$$\mathcal{F}_-(\mathbf{B}) = d\mathcal{A}_-(\mathbf{B}) = -\frac{i}{2} \frac{B_z dB_x \wedge dB_y + B_y dB_z \wedge dB_x + B_x dB_y \wedge dB_z}{B^3}. \quad (7.56)$$

The detailed calculations can be found in App. A.9. One can rewrite Eqs. (7.55) and (7.56) as

$$\mathcal{A}_-(\mathbf{B}) = -\frac{i}{2B^2} \frac{[\mathbf{e}_z \times \mathbf{B}]}{(\mathbf{e}_z \cdot \mathbf{B}/B - 1)} \begin{pmatrix} dB_x \\ dB_y \\ dB_z \end{pmatrix} =: \mathbf{A}_-(\mathbf{B}) d\mathbf{B}, \quad (7.57)$$

and

$$\mathcal{F}_-(\mathbf{B}) = -\frac{i\mathbf{B}}{B^3} \begin{pmatrix} dB_y \wedge dB_z \\ dB_z \wedge dB_x \\ dB_x \wedge dB_y \end{pmatrix} =: \mathbf{F}_-(\mathbf{B}) d\mathbf{S}. \quad (7.58)$$

Note that we have

$$\mathbf{F}_-(\mathbf{B}) = \text{curl} \mathbf{A}_-(\mathbf{B}) = \nabla \times \mathbf{A}_-(\mathbf{B}) \quad (7.59)$$

by definition, cf. App. A.9. Moreover, the coefficient field $\mathbf{F}_-(\mathbf{B})$ of the Berry curvature in Eq. (7.58) readily resembles the magnetic field

$$\vec{B}(\mathbf{B}) = \frac{\mu_0 q_{\text{mag}}}{4\pi} \frac{\mathbf{B}}{|\mathbf{B}|^3}, \quad (7.60)$$

of a magnetic point charge, $\rho_{\text{mag}} = q_{\text{mag}} \delta(\mathbf{B})$, situated at the origin $\mathbf{B} = 0$ in $\mathbf{B}$ -space. The spin-Chern number of the $U(1)$ subbundle $\psi_- \xrightarrow{\pi_-} S$ over a closed two-dimensional submanifold $S \subset \mathbb{R}^3 \setminus \{\mathbf{0}\}$ is given by the total spin-Berry flux

$$Ch_1^{(S)}(\mathcal{F}_-(\mathbf{B})) = \frac{i}{2\pi} \oint_S \mathcal{F}_-(\mathbf{B}) = \frac{1}{2\pi} \oint_S \frac{\mathbf{B}}{B^3} d\mathbf{S} \quad (7.61)$$

through S . In particular, we find that any surface S that encloses the magnetic charge at the origin gives (for details see App. A.9)

$$Ch_1^{(S)}(\mathcal{F}_-(\mathbf{B})) = 1. \quad (7.62)$$

The magnetic monopole analogy will be helpful below. Note that we largely follow the mathematical convention, treating the Berry connection $\mathcal{A}_-(\mathbf{B})$ as a one-form and the Berry curvature $\mathcal{F}_-(\mathbf{B})$ as a two-form. In contrast, the physics literature often focuses on the associated coefficient fields $\mathbf{A}_-(\mathbf{B})$ and $\mathbf{F}_-(\mathbf{B})$, treating them as vector fields over $\mathbf{B}$ -space. Further aspects of how the physics notation can be reconciled with the underlying mathematical concepts are addressed in App. A.9.

In the $J \rightarrow \infty$ limit, the rest of the system, i.e. the “punctured” Haldane model without the impurity site i_0 , does not couple to the impurity-spin manifold $\mathcal{S}_1$ at all, so it carries no spin-Chern number. For a single impurity spin in the limit $J \rightarrow \infty$, we may therefore analytically compute the spin-Chern number of the *full* model Eq. (7.1) using the *simplified monopole* model described above. This yields a total spin-Chern number of $Ch_1^{(S)} = 1$. Since the full model has $Ch_1^{(S)} = 0$ at $J = 0$ and $Ch_1^{(S)} = 1$ for $J \rightarrow \infty$, there must be a topological transition at some intermediate $J = J_{\text{crit}}$. This change is driven by a localised impurity and necessarily involves a gap closure caused by an in-gap state crossing the chemical potential.

7.5 Topological Transition at J_{crit}

The simple monopole model from Eq. (7.49) yields $Ch_1^{(S)} = 1$ for every $J > 0$, and an ill-defined spin-Chern number at $J = 0$, where the ground state of H_{mono} becomes twofold degenerate. This shows that the topological transition occurs exactly at $J = 0$. Intuitively, the spin-Chern number becomes ill-defined when the integration contour itself contains magnetic charge. This interpretation locates the magnetic monopole charge at $\mathbf{B}_{\text{crit}} = J\mathbf{S}/2 = \mathbf{0}$, precisely as built into the model from the start.

In contrast, the S -space topological phase transition of the *full* model in Eq. (7.1) occurs at a *finite* critical exchange coupling $J_{\text{crit}} > 0$. As a consequence, the accompanying gap closure is not constrained to a single $\mathbf{B}_{\text{crit}} \in \mathbb{R}^3 \setminus \{\mathbf{0}\}$ either. Instead, it takes place across an entire *critical surface*

$$S_{\text{crit}} = \{\mathbf{B} \in \mathbb{R}^3 \setminus \{\mathbf{0}\} \mid B = |\mathbf{B}| = J_{\text{crit}}S/2\}, \quad (7.63)$$

defined solely by the critical exchange interaction strength J_{crit} . The accompanying *infinite degeneracy* of the ground state energy at J_{crit} is caused by the aforementioned $\text{SO}(3)$ rotation symmetry of the total Hamiltonian in Eq. (7.1): at $J = J_{\text{crit}}$, we obtain two degenerate many-body ground states of the form given in Eq. (7.44) for *each fixed* $\mathbf{S} \in \mathbb{S}^2$. For the monopole analogy [81, 137], this implies that the magnetic charge q_{mag} is uniformly distributed over the critical *two-sphere* $\mathbb{S}_{\text{crit}}^2 \subset \mathbb{R}^3 \setminus \{\mathbf{0}\}$ with radius $B_{\text{crit}} = J_{\text{crit}}S/2$, i.e. we have a generalised monopole charge density

$$\rho_{\text{mag}}(\mathbf{B}) = \sigma_{\text{mag}} \delta(B - J_{\text{crit}}S/2) \quad (7.64)$$

with the magnetic surface charge density

$$\sigma_{\text{mag}} = \frac{q_{\text{mag}}}{4\pi J_{\text{crit}}^2 S^2}. \quad (7.65)$$

Solving $\text{div} \mathbf{F}(\mathbf{B}) = \mu_0 \rho_{\text{mag}}(\mathbf{B})$ using the divergence theorem and exploiting the $\text{SO}(3)$ symmetry yields the magnetic field

$$\vec{B}(\mathbf{B}) = \frac{\mu_0 q_{\text{mag}}}{4\pi} \frac{\mathbf{B}}{|\mathbf{B}|^3} \Theta(B - J_{\text{crit}}S/2), \quad (7.66)$$

and hence the Berry curvature coefficient field

$$\mathbf{F}_-(\mathbf{B}) = -\frac{i\mathbf{B}}{B^3} \Theta(B - J_{\text{crit}}S/2). \quad (7.67)$$

Here, Θ is the Heaviside step function. Due to $\mathbf{B} = J\mathbf{S}/2$, it follows that $\Theta(B - J_{\text{crit}}S/2) = \Theta(J - J_{\text{crit}})$. Accordingly, the curvature coefficient field $\mathbf{F}_-(\mathbf{B})$ vanishes inside the critical sphere, $J < J_{\text{crit}}$, while it coincides with the field of the magnetic point charge outside the critical sphere, i.e. for $J > J_{\text{crit}}$.

As before, the spin-Chern number characterising the subbundle $\psi_- \xrightarrow{\pi_-} \mathbb{S}_{JS/2}^2$ over a two-sphere $\mathbb{S}_{JS/2}^2$ with radius $B = JS/2$ is defined as

$$Ch_1^{(S)}(\mathcal{F}_-(\mathbf{B})) = \frac{1}{2\pi} \oint_{\mathbb{S}_{JS/2}^2} \mathbf{F}_-(\mathbf{B}) d\mathbf{S} = \Theta(J - J_{\text{crit}}), \quad (7.68)$$

where $\mathcal{F}_-(\mathbf{B})$ is now the Berry curvature two-form associated with the coefficient field $\mathbf{F}_-(\mathbf{B})$ from Eq. (7.66). Interestingly, the Chern number in Eq. (7.68) jumps from $Ch_1^{(S)} = 0$ for $J < J_{\text{crit}}$ to $Ch_1^{(S)} = 1$ for $J > J_{\text{crit}}$. The connection to the curvature from Eq. (7.67) is once more given by [2]

$$\mathbf{A}(\mathbf{B}) = -\frac{i}{2B^2} \frac{[\mathbf{e}_z \times \mathbf{B}]}{(\mathbf{e}_z \cdot \mathbf{B}/B - 1)} \Theta(J - J_{\text{crit}}). \quad (7.69)$$

The z -unit vector $\mathbf{e}_z$ in Eq. (7.69) determines the direction of the Dirac string singularity, which, in this case, appears along the negative z -axis for $|z| > J_{\text{crit}}S/2$. While this singularity can be displaced by gauge transformations $\mathbf{A}(\mathbf{B}) \mapsto \mathbf{A}(\mathbf{B}) + \nabla \Lambda(\mathbf{B})$ with an arbitrary scalar field Λ , it cannot be eliminated entirely. This accounts for the fact that a non-trivial first Chern number $Ch_1^{(S)} \neq 0$ obstructs the definition of a nowhere vanishing global section, cf. e.g. the discussion following Thm. 2.3.2 in Sec. 2.3 and App. A.9. Inside the critical sphere we find $\mathbf{F}(\mathbf{B}) = \text{rot} \mathbf{A}(\mathbf{B}) = 0$, so there exists a local gauge choice in which $\mathbf{A}(\mathbf{B}) = 0$. The Berry connection is discontinuous on the critical sphere and along the Dirac string stretching from the point $-J_{\text{crit}}S/2 \mathbf{e}_z$ on the critical sphere to infinity.

7.6 Relation to k -Space Topology

The topological transition at J_{crit} is driven by the electronic structure around the impurity spin $\mathbf{S}$. As we have seen earlier, $J\mathbf{S}$ acts as a local magnetic field, which polarises the local magnetic moment of the electron system and results in the formation of two high-energy Zeeman states $|\varepsilon_{i_0\downarrow}\rangle$ and $|\varepsilon_{i_0\uparrow}\rangle$. Irrespective of the other parameters of the electronic model, the spin-down (spin-up) state moves downwards (upwards) in energy as J increases, and eventually separates from the lower (upper) edge of the valence (conduction) band at a coupling strength roughly set by the width of the valence band $J \sim W_{\text{val}}$ (conduction band $J \sim W_{\text{cond}}$). Note that the energies of these two states are not visible in Fig. 7.3. Both high-energy states are exponentially localised in the vicinity of the impurity site i_0 and become fully localised at i_0 for $J \rightarrow \infty$, at which point they realise the magnetic-monopole model from Eq. (7.49).

The physical origin of the low-energy states *within* the bulk band gap is more subtle. We find that these states are strongly influenced by the k -space topological phase of the electron system, characterised by the first k -space Chern number $C_1^{(k)}$. In the k -space topologically non-trivial electron system with $C_1^{(k)} = \pm 1$, we always obtain *two* in-gap states for sufficiently strong J . These states remain isolated from the bulk continuum and stay within the gap for all J , cf. states (a), (b) and (c), (d) in Fig. 7.3. As $J \rightarrow \infty$, their energies eventually become degenerate. For $C_1^{(k)} = 0$, on the other hand, there is only *one* in-gap state, which fully crosses the gap as function of J . In particular, this state merges with the conduction or valence-band continuum at some finite J . As a result, there is no in-gap state present in the $J \rightarrow \infty$ limit in this case, cf. state (e) in the right panel of Fig. 7.3.

For $J \rightarrow \infty$, the classical impurity spin coupled to i_0 becomes a hard zero-dimensional defect in both spin sectors of the Haldane model. According to the tenfold way, codimension-two defects in two-dimensional Altland-Zirnbauer class A systems are topologically trivial [14, 139]. Accordingly, the bulk-defect correspondence does not predict a topologically protected defect mode localised at i_0 in the non-trivial $C_1^{(k)} = \pm 1$ phase. On the other hand, for a soft defect with finite impurity strength, it is well-known that impurity bound states *can* serve as a local signature of the bulk topological phase. This has, for instance, been demonstrated explicitly for codimension-two defects in two-dimensional $\mathbb{Z}_2$ insulators [126, 131]. In the present case of a magnetic point impurity in a Chern insulator, one might expect a strong connection between the k -space Chern number and the existence of in-gap impurity bound states as well. Indeed, such states have been observed for the Haldane model with various types of spinless local impurity potentials [135].

Here, we present numerical evidence and an intuitive argument that, in the strong- J limit, a localised spin-up and spin-down in-gap state must emerge *if and only if* the electronic bulk structure is topologically non-trivial. To see this, recall that for $J \rightarrow \infty$, the impurity-bound spin-down state $|\varepsilon_{i_0\downarrow}\rangle$ becomes fully occupied, while its spin-up partner $|\varepsilon_{i_0\uparrow}\rangle$ is pushed so far up in energy that it is effectively removed. As a result, hopping to and from the impurity site is suppressed, and the electronic structure of Eq. (7.1) becomes equivalent to that of a spinful Haldane model with a single-site sized *hole* at i_0 . Let us imagine that we can *inflate* this hole up to macroscopic size, cf. Ref. [129]. The edge of the resulting macroscopic hole constitutes a one-dimensional defect in the periodic Haldane model, which, due to the conventional bulk-boundary correspondence, supports one chiral mode per spin projection if $C_1^{(k)} \neq 0$. This mode traverses the bulk band gap and disperses with quasi-momentum parallel to the edge of the hole. For a hole with a finite circumference, the quasi-momentum along its edge becomes discretised, leading to a finite number of N_{edge} in-gap states localised along the edge. Here, $N_{\text{edge}} = q C_{\text{edge}}$ is a fraction $q \in \mathbb{Q}$ of the total number of edge sites C_{edge} around the circumference of the hole. Consequently, shrinking the large hole back to its original single-site size gradually coarsens the quasi-momentum discretisation, decreasing the number of in-gap edge modes until only a single impurity-bound state per spin projection remains. It is worth mentioning that for $J \rightarrow \infty$, the two super-discretised bound states associated with the two spin projections have the same energy because the only available mechanism capable of lifting their degeneracy – coupling to the impurity spin – affects only the impurity site in that case. For finite but large J , we may invoke a perturbative argument: the correction to the bound-state energies due to second-order virtual-hopping processes to the impurity site i_0 and back is of the order of t_{hop}^2/J . Indeed, the energies of the in-gap states (a), (b) and of (c), (d) in Fig. 7.3 are approximately proportional to $1/J$ for strong J , and the correction is negative (positive) for spin-up (spin-down) states.

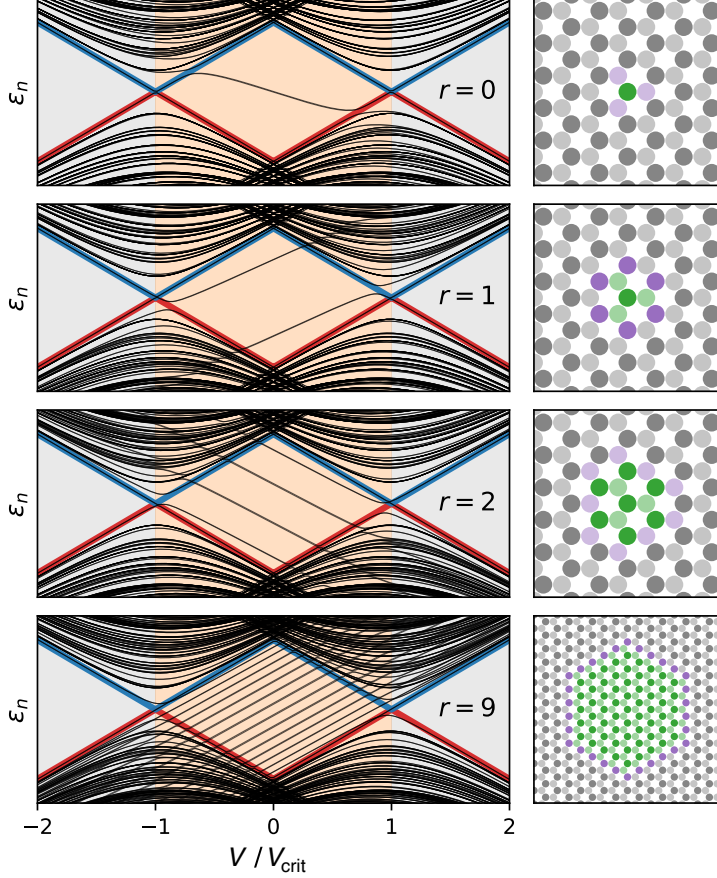


Figure 7.4: Single-particle energies as a function of sublattice potential V (left panels) for the Haldane model with holes of different radii $r = 0, 1, 2, 9$ (from top to bottom), all centred on site i_0 . See main text for the definition of r . The holes (right panels) are created by deleting all hoppings between the hole sites (green) and the surrounding lattice, effectively removing them. Violet sites mark the first shell outside the hole. Dark and light grey colours indicate sites of sublattice A and B sublattices, respectively. Parameters as in Fig. 7.3. The straight red and blue lines represent the V -dependent bulk band gap. All energies ε_n are twofold spin degenerate. Adapted with minor modifications from Ref. [RQ1].

The previously described coarsening of the in-gap edge-bound states is demonstrated in Fig. 7.4, which shows the discrete low-energy spectrum of the electron system as a function of V/V_{crit} for periodic Haldane models with holes of three different sizes. The holes are created by removing clusters

$$C_r = \{j \in \Lambda_h \mid d_h(i_0, j) \leq r\} \quad (7.70)$$

of lattice sites $j \in \Lambda_h$, whose graph distance $d_h(i_0, j)$ to a distinguished site i_0 does not exceed r . Here, the honeycomb lattice Λ_h is treated as an undirected graph $\Lambda_h = (V_h, E_h)$, in which the vertex set V_h corresponds to the lattice sites and the edge set E_h contains unordered pairs of nearest-neighbour sites (NN bonds). In this description, a path from i to j is defined as a finite subgraph $\gamma = (V_\gamma, E_\gamma) \subset \Lambda_h$ with vertex set

$$V_\gamma = \{v_1 = i, v_2, \dots, v_n = j\} \subset V_h \quad (7.71)$$

of adjacent lattice sites, and edge set

$$E_\gamma = \{(v_1, v_2), (v_2, v_3), \dots, (v_{n-1}, v_n)\} \subset E_h \quad (7.72)$$

of pairwise NN bonds connecting them. The length $|\gamma|$ of a path γ is given by the number of edges it contains, i.e. $|\gamma| = |E_\gamma|$. If we denote the set of all paths from site i to site j by Γ_{ij} , we may define the graph distance as a function

$$d_h : \Lambda_h \times \Lambda_h \rightarrow \mathbb{N}_0, \quad (i, j) \mapsto \min_{\gamma \in \Gamma_{ij}} (|\gamma|), \quad (7.73)$$

which assigns to each pair (i, j) of lattice sites (vertices) the length of the shortest path connecting them. The function d_h defines a graph metric on Λ_h that we call the honeycomb metric. With this, the parameter r in Eq. (7.70) determines the radius of a spherical hole in the honeycomb metric d_h .

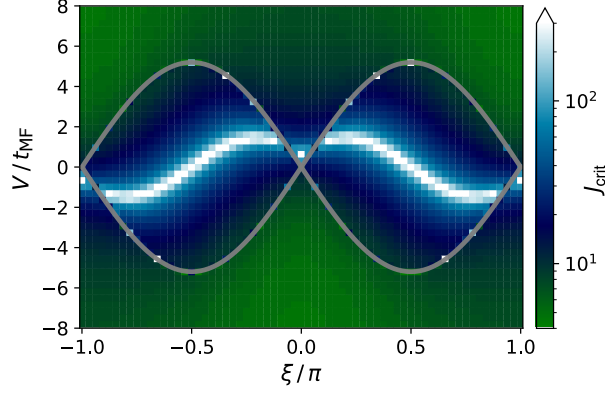


Figure 7.5: Critical interaction strengths J_{crit} (colour bar) marking the S -space topological transition between the $Ch_1^{(S)} = 0$ phase $J < J_{\text{crit}}$ and the $Ch_1^{(S)} = 1$ phase for $J > J_{\text{crit}}$, as a function of V and ξ . Results are for a single classical spin ($R = 1$); other parameters match those in Fig. 7.3. Adapted with minor modifications from Ref. [RQ1].

For the largest considered hole radius of $r = 9$ (bottom panel), the hole consists of 136 removed sites (green dots), while the first outer shell (violet dots) at distance $r = 10$ from i_0 is formed by $3r = 30$ sites of sublattice A (dark dots). We see that for any sublattice potential with $-V_{\text{crit}} < V < V_{\text{crit}}$ (light orange in Fig. 7.4) that puts the Haldane model in the topologically non-trivial phase, there are states inside the V -dependent band gap at nearly equidistant energies. This equidistance reflects the almost linear dispersion of the boundary modes in the Haldane model. In the $r \rightarrow \infty$ limit, the in-gap states would completely fill the band gap. For smaller r , e.g. $r = 2$ (second panel from the bottom), the spectral flow with V shows no qualitative change, except that the number of in-gap states is reduced due to the smaller number of edge sites.

The in-gap states are exponentially localised at the edge, i.e. on the first outer shell (violet sites). Beyond this shell, their weight decays rapidly with increasing r . The bipartiteness of the honeycomb lattice ensures that the edge sites belong exclusively to sublattice A (B) when r is even (odd). Accordingly, the energies of the in-gap states are expected to increase (decrease) linearly with V for even (odd) r . This is nicely confirmed by the numerical calculations shown in Fig. 7.4.

By gradually shrinking the hole, we eventually get a single (spin-degenerate) impurity mode that is exponentially localised on the three nearest neighbours of the impurity site i_0 , cf. top panel of Fig. 7.4. This corresponds to the in-gap mode shown in Fig. 7.3 (left and middle) for strong J , where it is slightly spin-split, see states (a), (b) and (c), (d). As described, this state is a super-discretised remnant of a topologically protected chiral mode localised on the one-dimensional boundary of a hypothetical macroscopic hole. In this sense, it is rooted in the k -space topological state of the host system and in the bulk-boundary correspondence for a codimension-one defect.

On the other hand, its existence cannot be fully explained within the tenfold way classification alone. A topologically protected defect state, as enforced by the bulk-defect correspondence, would be pinned to the chemical potential μ . Yet, as seen in Fig. 7.4 (top), the energy of the $r = 0$ mode varies with V , contradicting an interpretation in terms of a codimension-two topological mode. While this behaviour *is* fully consistent with the tenfold way – which does not predict topological zero modes at point-like defects in the symmetry class (A) of the Haldane model [14, 139] – it highlights the need for an alternative explanation for the emergence of this localised in-gap state.

To further clarify the role of k -space topology, we briefly examine the case with $C_1^{(k)} = 0$, where the k -space topology is trivial. In this regime, no in-gap state appears around i_0 for strong J , see the ranges $V < -V_{\text{crit}}$ and $V > V_{\text{crit}}$ in Fig. 7.4. This is consistent with the lack of a dispersive edge mode at the one-dimensional boundary of a trivial Chern insulator. However, the change of the spin-Chern number from $Ch_1^{(S)} = 0$ at $J = 0$ to $Ch_1^{(S)} = 1$ at $J = \infty$ continues to enforce the appearance of an in-gap state in some *intermediate* coupling-strength range. This state must bridge the band gap as function of J , see state (e) in Fig. 7.3 (right). Although this mode is also localised in the vicinity of i_0 it has much less weight close to i_0 than the high-energy bound states.

The critical interaction J_{crit} , at which the spin-Chern number jumps from $Ch_1^{(S)} = 0$ to $Ch_1^{(S)} = 1$, is determined by the electronic bulk of the host system. Figure 7.5 shows numerical results for the critical coupling strength J_{crit} as a function of the model parameters ξ and V for fixed NNN hopping strength of $t_{\text{MF}} = 0.1$. The asymmetry of the phase diagram with respect to V arises from the fact that the impurity is coupled to a site on sublattice A. Coupling it to a site of sublattice B would reverse the role of $V \rightarrow -V$ and mirror the phase diagram about the $V = 0$ axis.

Generally, the local topological transition to a finite spin-Chern number requires a strong exchange coupling, typically on the order of the band width or stronger. Based on the previous discussion, one would expect J_{crit} to be larger when the host system is in a k -space topologically non-trivial phase, as in this case the spin-split in-gap states remain within the gap for $J \rightarrow \infty$. By contrast, the in-gap states of the k -space topologically trivial phase only appear in an intermediate region, where they “quickly” bridge the bulk gap and vanish into the bulk states for greater coupling strengths.

This expectation is corroborated by the results shown in Fig. 7.5, where the k -space topological phase-transition line of the pristine Haldane model is indicated by the thick grey lines (see also Fig. 7.2). The critical coupling strength J_{crit} in the non-trivial Chern insulating phase is roughly an order of magnitude larger than in the trivial phase. For certain parameters, J_{crit} appears to diverge, see the white curve in Fig. 7.5. On this curve, the Zeeman pair of spin-up and spin-down in-gap states that emerges for strong J in the k -space topological phase is located symmetrically around the chemical potential, so that neither state crosses μ as a function of J . Consider, for instance, $\xi = \pi/2$ and $V = 0$. There, particle-hole symmetry requires $\mu = 0$ and a symmetric spin splitting of the in-gap states around μ for all J , implying $J_{\text{crit}} = \infty$. For $\xi \neq \pi/2$, there is a unique $V < \infty$, for which the in-gap states never cross μ .

7.7 Two Impurity Spins

For $R = 2$ classical spins, the configuration space manifold becomes $\mathcal{S}_2 = \mathbb{S}^2 \times \mathbb{S}^2$, which is a simply-connected, closed, compact and orientable manifold of real dimension $\dim_{\mathbb{R}}(\mathcal{S}_2) = 4$. As before, we use coordinates $\boldsymbol{\lambda} = (\lambda_0, \lambda_2, \lambda_2, \lambda_3) \equiv (\vartheta_0, \phi_0, \vartheta_1, \phi_1)$ on $\mathcal{S}_2$, where each pair (ϑ_j, ϕ_j) denotes the polar and azimuthal angles on the j -th factor $\mathbb{S}_j^2$ of the product manifold $\mathcal{S}_R$. The second spin-Chern number $Ch_2^{(S)}$ is obtained from Eq. (7.9). Once more, $Ch_2^{(S)}$ must vanish for $J = 0$, since the manifold of spin configurations $\mathcal{S}_2$ is completely decoupled from the electron degrees of freedom in that case. Similarly, the local physics at the two sites i_0 and i_1 , to which $\mathbf{S}_0$ and $\mathbf{S}_1$ are coupled, is again captured by

$$H_{2\text{-mono}} = J\mathbf{S}_0\mathbf{s}_{i_0} + J\mathbf{S}_1\mathbf{s}_{i_1}. \quad (7.74)$$

for $J \rightarrow \infty$. In this limit, the ambient system becomes a doubly-punctured Haldane model with holes at i_0 and i_1 . Since this system is not coupled to $\mathcal{S}_2$, it carries no spin-Chern number. Meanwhile, the second spin-Chern number of the two isolated magnetic monopoles in Eq. (7.74) is simply the product

$$Ch_2^{(S)} = Ch_1^{(S)} \cdot Ch_1^{(S)} = 1 \quad (7.75)$$

of the two respective first spin-Chern numbers associated with the isolated monopoles. For $J \rightarrow \infty$, we therefore have a second spin-Chern number of

$$Ch_2^{(S)} = 1 \quad (7.76)$$

for the entire quantum-classical hybrid system. As before, this indicates that there must be a transition between two topologically different local phases as a function of J . Note that the same factorisation of the second spin-Chern number occurs at finite J provided the impurity sites i_0 and i_1 are sufficiently far apart, in which case the two-impurity problem decouples into two separate single-impurity problems regardless of the explicit strength of J .

For the numerical calculations, we use the same model parameters as in the single-spin ($R = 1$) case; see caption of Fig. 7.3. The two impurity spins are coupled to NNN sites i_0 and i_1 of sublattice A. Due to the $\text{SO}(3)$ symmetry of the Hamiltonian, the single-particle spectrum $\varepsilon_n(\mathbf{S}_0, \mathbf{S}_1)$ depends only on the relative orientations of $\mathbf{S}_1$ and $\mathbf{S}_2$ captured by their scalar product $\mathbf{S}_0\mathbf{S}_1 = \cos(\varphi)$, where φ denotes the relative angle between the impurity spins $\mathbf{S}_0$ and $\mathbf{S}_1$. This allows us to write $\varepsilon_n(\mathbf{S}_0, \mathbf{S}_1) = \varepsilon_n(\varphi)$.

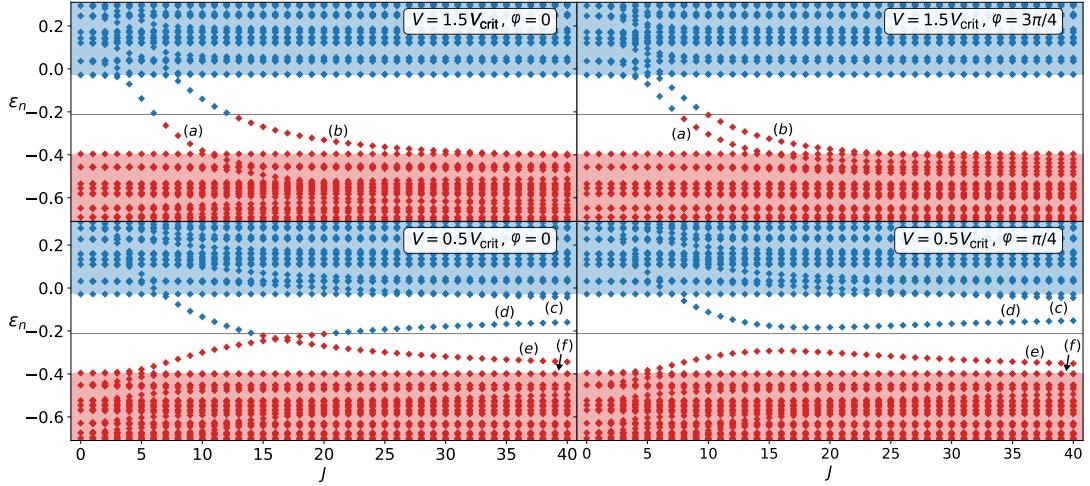


Figure 7.6: Low-energy single-particle spectrum as a function of J . Analogous to Fig. 7.3, but for two impurity spins $\mathbf{S}_0$ and $\mathbf{S}_1$ coupled to next-nearest-neighbour sites belonging to sublattice A of the honeycomb lattice. Calculations performed for various on-site potentials V and angles φ between $\mathbf{S}_0$ and $\mathbf{S}_1$ as indicated. Further parameters are $t_{\text{hop}} = 1$, $\xi = \pi/4$, $t_{\text{MF}} = 0.1$, 39×39 unit cells, and $\mu \approx -0.21$ in the middle of the bulk band gap (gray line). In-gap bound states are labelled by (a) – (e). State (f) is part of the bulk continuum. In the k -space trivial phase, states (a) and (b) cross the gap within a finite J range. Their energy splitting decreases with increasing φ and vanishes for $\varphi = \pi$. In the k -space non-trivial phase, states (c) and (d), as well as states (e) and (f), form Zeeman-split pairs whose energies become degenerate for $J \rightarrow \infty$. Adapted with minor modifications from Ref. [RQ1].

Figure 7.6 shows the low-energy spectrum of single-particle energies ε_n around $\mu \approx -0.21$ as a function of J for two different sublattice potentials V . The top panels present results for $V = 1.5V_{\text{crit}}$ (k -space topologically trivial) with spin configurations $\varphi = 0$ (left) and $\varphi = 3\pi/4$ (right), while the bottom panels correspond to $V = 0.5V_{\text{crit}}$ (k -space topologically non-trivial) with $\varphi = 0$ (left) and $\varphi = \pi/4$ (right).

We first discuss the k -space topologically trivial phase of the host system where $C_1^{(k)} = 0$, shown in the upper panels. In contrast to the single-impurity case ($R = 1$), the system now hosts *two* in-gap states that fully bridge the bulk band gap as a function of J , see states (a) and (b) in Fig. 7.6. For $\varphi = 0$, these states cross the chemical potential at critical couplings

$$J_1(\varphi = 0) \approx 6.0 \quad \text{and} \quad J_2(\varphi = 0) \approx 12.3. \quad (7.77)$$

As φ increases (upper right panel of Fig. 7.6 for $\varphi = 3\pi/4$), the critical coupling $J_1(\varphi)$ increases, while $J_2(\varphi)$ decreases until they coincide at $\varphi = \pi$, where $J_1(\varphi = \pi) = J_2(\varphi = \pi)$.

The equality of $J_1(\varphi)$ and $J_2(\varphi)$ at $\varphi = \pi$ can be understood as follows. At $\varphi = \pi$ the impurity spins are precisely antiparallel. As a result, the z -component $s_{\text{tot},z}$ of the total electron spin is conserved, and the two impurity bound states have well-defined and opposite spin-projection quantum numbers. This prevents hybridisation among the bound states. Moreover, since the states are related by a symmetry transformation – specifically, a spin flip ($\uparrow \longleftrightarrow \downarrow$) combined with a reflection about the bond-centered mirror axis passing through the shared NN site on sublattice B – their energies must be degenerate for all J . Consequently, they cross the chemical potential μ at the same critical coupling strength $J_1(\varphi = \pi) = J_2(\varphi = \pi)$.

While the *antiparallel* impurity-spin configuration ($\varphi = \pi$) corresponds to *minimal* hybridisation, and hence minimal $\Delta J(\varphi = \pi) \equiv J_2(\varphi = \pi) - J_1(\varphi = \pi) = 0$, the *parallel* configuration ($\varphi = 0$) yields *maximal* hybridisation, resulting in the maximal difference $\Delta J(\varphi = 0) = \max_{\varphi} \{\Delta J(\varphi)\}$ of critical coupling strengths. For intermediate angles φ with $0 < |\varphi| < \pi$, the hybridisation among bound states and the critical coupling difference $\Delta J_{\text{crit}}(\varphi)$ vary continuously between these extremes. This is illustrated in the left panel of Fig. 7.7 for the k -space topologically trivial system. We find that for every $J_{\text{crit},1} < J < J_{\text{crit},2}$ there exists exactly one $0 < \varphi < \pi$, for which the system becomes gapless, see boundaries between areas of different colours in the left panel of Fig. 7.7. Accordingly, the overall critical couplings

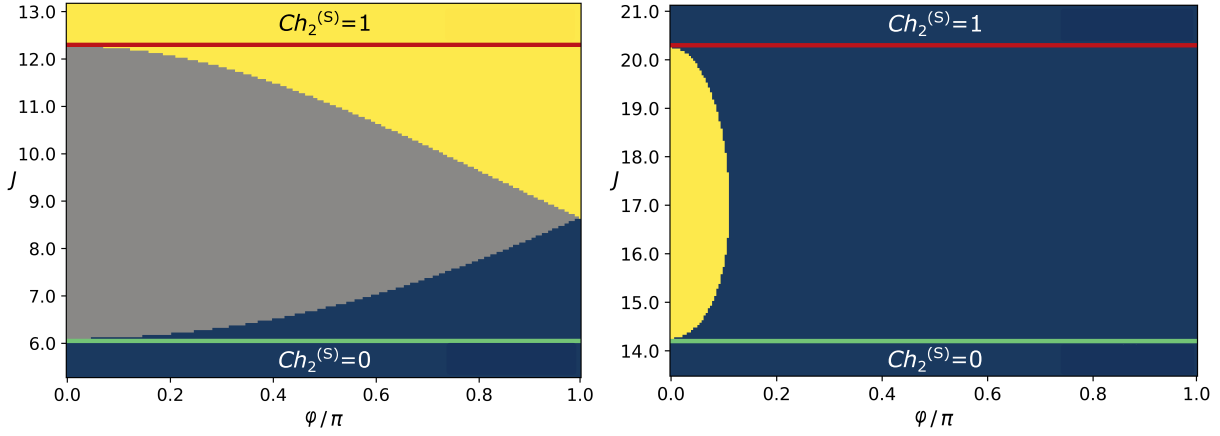


Figure 7.7: Transition of the second spin-Chern number $Ch_2^{(S)}$ as a function of the exchange coupling strength J . Boundaries between coloured areas correspond to critical exchange couplings, $J_1(\varphi)$ and $J_2(\varphi)$, at which an in-gap state crosses μ as function of the angle φ enclosed by $\mathbf{S}_0$ and $\mathbf{S}_1$. The host system is a spinful Haldane model with $t_{MF} = 0.1$ and $\xi = \pi/4$, defined on a periodic honeycomb lattice of 27×27 unit cells. The two impurities couple to next-nearest neighbour sites of sublattice A. *Left:* $V/V_{\text{crit}} = 1.5$ (k -space trivial phase), $Ch_2^{(S)} = 0$ for $J < J_{\text{crit},1} \approx 6.0$ (green line), $Ch_2^{(S)} = 1$ for $J > J_{\text{crit},2} \approx 12.3$ (red line). *Right:* $V/V_{\text{crit}} = 0.5$ (k -space non-trivial phase), $Ch_2^{(S)} = 0$ for $J < J_{\text{crit},1} \approx 14.2$ (green line), $Ch_2^{(S)} = 1$ for $J > J_{\text{crit},2} \approx 20.3$ (red line). Adapted with minor modifications from Ref. [RQ1].

are given by

$$J_{\text{crit},1} \equiv J_1(\varphi = 0) \approx 6.0 \quad \text{and} \quad J_{\text{crit},2} \equiv J_2(\varphi = 0) \approx 12.3. \quad (7.78)$$

Note that the second spin-Chern number is ill-defined for every J in this critical range. In this sense, the critical range separates the trivial phase at $J < J_{\text{crit},1} \approx 6.0$ with $Ch_2^{(S)} = 0$ from the non-trivial phase at $J > J_{\text{crit},2} \approx 12.3$ with $Ch_2^{(S)} = 1$. Notably, the critical interval $[J_{\text{crit},1}, J_{\text{crit},2}]$ of the two-impurity ($R = 2$) system is roughly centred around the critical coupling J_{crit} of the single-impurity ($R = 1$) case; as seen by comparing the top panels of Fig. 7.6 to the right panel of Fig. 7.3. This can be understood in the limit of large impurity separation: as the distance between i_0 and i_1 is increased, the electronic environments around the two spins decouple, the in-gap states become degenerate, and the critical coupling becomes independent of the relative angle φ .

We now turn to the k -space topologically non-trivial case where $C_1^{(k)} = 1$, shown in the bottom panels of Fig. 7.6. In the limit of infinite impurity separation, the strong- J regime features four in-gap states: two states localised around i_0 and two states localised around i_1 . Both state pairs constitute super-discretised remnants of topologically protected chiral modes localised at the edges of two perfectly separated macroscopic holes centered at i_0 and i_1 , as described above. The energies of the in-gap states localised around different i_0 and i_1 are degenerate in this limit. With decreasing inter-impurity distance, and increasing overlap between the in-gap states, their degeneracy is lifted by the formation of bonding and anti-bonding superpositions. This results in two Zeeman-split pairs of in-gap states that appear at distinct energies in the bulk gap for strong but finite coupling strengths J . At the minimal non-trivial inter-impurity distance, i.e. when i_0 and i_1 are nearest-neighbour sites of opposite sublattices, they effectively combine into a single two-site hole. As illustrated in Fig. 7.4, a single four-site hole ($r = 1$) at $V/V_{\text{crit}} = 0.5$ yields a single pair of in-gap states. By analogy, one expects a single pair of in-gap states in case of a NN two-site hole, too. This indicates that decreasing the inter-impurity distance, causes one pair of in-gap states to merge with the continuum of delocalised bulk states. The configuration, in which the impurity sites i_0 and i_1 occupy next-nearest-neighbour sites of the same sublattice, constitutes an *intermediate* case between the adjacent (single two-site hole) and the infinitely separated (two single-site holes) impurity limits. As shown in the bottom panels of Fig. 7.6, there are three in-gap states for large but finite J . For both spin configurations, $\varphi = 0$ and $\varphi = \pi/4$, the energies of the Zeeman pair formed by

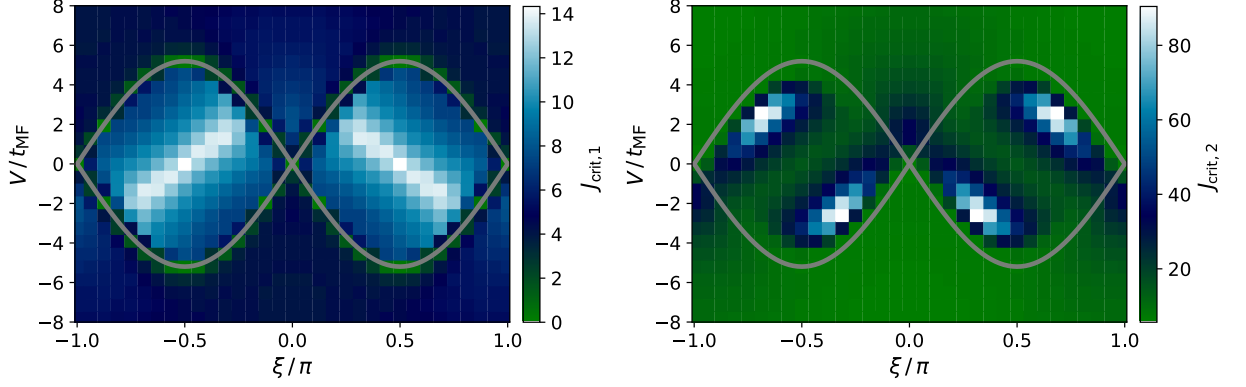


Figure 7.8: Functional dependence of the critical coupling strength on the model parameters V and ξ . Numerical results for a spinful Haldane model with $t_{\text{MF}} = 0.1$ defined on a periodic honeycomb lattice of 9×9 unit cells. The two impurities are coupled to next-nearest neighbour sites of sublattice A. The chemical potential is located in the middle of the bulk band gap. Thick gray lines mark the k -space topological phase boundaries of the Haldane model. *Left:* Lower critical interaction $J_{\text{crit},1}$ (colour bar) for the S -space topological transition from the trivial $Ch_2^{(S)} = 0$ phase at weak J to the gapless phase with undefined spin-Chern number. *Right:* Upper critical interaction $J_{\text{crit},2}$ (colour bar) for the S -space topological transition from the gapless phase at intermediate J to the non-trivial $Ch_2^{(S)} = 1$ phase at strong J . Note the different scales on the two colour bars. Adapted with minor modifications from Ref. [RQ1].

states (c) and (d) remain within the gap and become degenerate for $J \rightarrow \infty$. In the other Zeeman pair, only state (e) remains inside the gap, while state (f) stays part of the bulk continuum as $J \rightarrow \infty$. For weaker couplings $J \lesssim 30$ and for $\varphi = 0$, only the spin-up state (d) (spin-down state (e)) remains within the gap and decreases (increases) in energy with decreasing J . For parallel impurity-spin alignment ($\varphi = 0$), we observe a level crossing around $J \approx 16$. For the non-parallel configuration ($\varphi = \pi/4$), this crossing is avoided, indicating a finite hybridisation between the involved states. This hybridisation stems from the relative orientation of the local spin quantisation axes. In the parallel case, the impurity-induced bound states at $\mathbf{S}_0$ and $\mathbf{S}_1$ are fully spin-polarised along a common axis. Expressed in either local quantisation frame, the spin-up state at $\mathbf{S}_0$ and spin-down state at $\mathbf{S}_1$ remain orthogonal, leading to a level crossing without hybridisation. In contrast, any non-parallel configuration of $\mathbf{S}_0$ and $\mathbf{S}_1$ defines two different local quantisation axes: in the $\mathbf{S}_0$ quantisation frame, the spin-up and spin-down states at $\mathbf{S}_1$ become superpositions of both $\mathbf{S}_0$ spin projections, and vice versa. The resulting non-orthogonality of the spin-polarised bound states at $\mathbf{S}_0$ and $\mathbf{S}_1$ enables their mixing, leading to hybridisation and the observed avoided crossing. Just like in the k -space topologically trivial case, there must be a gap closure in a critical- J range on the spin-configuration manifold $\mathcal{S}_2$, supporting the transition region between the S -space topologically trivial phase with $Ch_2^{(S)} = 0$ at $J = 0$ to the S -space topologically non-trivial phase with spin-Chern number $Ch_2^{(S)} = 1$ at $J \rightarrow \infty$.

As shown in the right panel of Fig. 7.7, there is a gap closure for $J_{\text{crit},2} \equiv J_2(\varphi = 0) \approx 20.3$. When J is decreased past $J_{\text{crit},2}$, the gap closes at $J_2(\varphi) \leq J_2(0)$ for increasing angle $\varphi = \arccos(\mathbf{S}_0 \mathbf{S}_1)$ until the angle reaches $\varphi \approx 0.11\pi$. For even smaller J , as described by the function $J_1(\varphi)$, the gap closure on $\mathcal{S}_2$ moves back to $\varphi = 0$ at $J_{\text{crit},1} \equiv J_1(\varphi = 0) \approx 14.2$. Note that due to the $\text{SO}(3)$ symmetry and for a coupling strength $J_{\text{crit},1} < J < J_{\text{crit},2}$, the gap closes on the whole three-dimensional submanifold of $\mathcal{S}_2$ determined by a critical $\varphi_{\text{crit}} = \varphi(J)$. In summary, the system is gapless within the critical range $J_{\text{crit},1} < J < J_{\text{crit},2}$, given by

$$J_{\text{crit},1} \equiv J_1(\varphi = 0) \approx 14.2 \quad \text{and} \quad J_{\text{crit},2} \equiv J_2(\varphi = 0) \approx 20.3. \quad (7.79)$$

The S -space topologically trivial phase is realised for $J < J_{\text{crit},1}$, while the S -space topologically non-trivial phase is found for $J > J_{\text{crit},2}$.

Figure 7.8 shows the functional dependence of the critical couplings $J_{\text{crit},1}$ and $J_{\text{crit},2}$ on the Haldane model parameters ξ and V . Specifically, the left panel displays the parameter dependence of $J_{\text{crit},1}$, while the right panel shows that of $J_{\text{crit},2}$. Note that both diagrams use the same colour map, but apply it to substantially different data ranges. This is necessary because $J_{\text{crit},2}$ exhibits much greater variations than $J_{\text{crit},1}$. In particular, $J_{\text{crit},2}$ can diverge when the corresponding Zeeman pair of in-gap states approaches μ symmetrically for $J \rightarrow \infty$; compare to the discussion of the white curves in Fig. 7.5. The rescaled colour map in the right panel of Fig. 7.8 helps to resolve parameter configurations with diverging $J_{\text{crit},2} \rightarrow \infty$ (white areas) in greater detail.

Another consequence of the rescaled colour maps is that similar values of $J_{\text{crit},1}$ and $J_{\text{crit},2}$ appear in quite different colours. To facilitate a direct comparison between the two panels, we recall that there is a level crossing of in-gap states for $\vartheta = 0$. When this level crossing occurs precisely at the chemical potential μ , the critical couplings coincide, i.e.

$$J_{\text{crit},1} = J_{\text{crit},2} . \quad (7.80)$$

With μ located in the middle of the bulk band gap, this degeneracy of in-gap states arises when $V = 0$ and $\xi = \pm\pi/2$. The corresponding points in the diagrams of Fig. 7.8 appear in white (left panel) and dark green (right panel), providing a visual anchor between the two colour scales.

Consistent with the single-impurity case, we observe that the topological transition characterised by the second spin-Chern number $C_2^{(\text{S})}$ typically occurs at stronger exchange couplings $J_{\text{crit},1}$ and $J_{\text{crit},2}$ when the host system is in the k -space topologically non-trivial phase.

8 – Long-Range Helical Spin Control

Boundary states arising from topological bulk-boundary correspondence have a number of highly desirable properties: they bridge the insulating bulk gap at the boundary, are protected against (symmetry-preserving) perturbations, and often exhibit anomalous physical properties. A prominent example of this are the topological edge states of the quantum spin Hall effect (QSHE), which are protected by TRS and exhibit spin-momentum locking. Given this unique combination of features, it becomes natural to ask how they can be harnessed for practical applications.

One area where this is particularly relevant is topological quantum computation [140, 141], which utilises, for example, fractional quantum Hall states [101] or Majorana zero modes [142] of topological superconductors. In these systems, the topological protection and the associated robustness are decisive factors. In the field of spintronics [143], examples include one-dimensional spin transport in inverted-gap semiconductor-based devices [144], and spin-dependent reflection with control of the spin rotation in trilayer junctions consisting of QSH and metallic materials [145]. The QSHE can be utilised to create nearly fully spin-polarised charge currents, controlled via magnetic defects [146], and fully electrical routes have been suggested to manipulate the spin of a magnetic adatom at the edge of a QSH insulator [147, 148].

Another approach is to focus on the interaction between the topological boundary states and suitable perturbations. A particularly promising type of perturbations are local impurities that intentionally break the protective symmetry [149]. This symmetry breaking allows the edge modes to interact non-trivially with the perturbation [150, 151], while the locality ensures that the topological protection and properties are preserved away from it. By exciting the boundary modes in a way that exploits their unique transport properties – such as spin-momentum locking in the QSHE or chirality in the QHE – one may establish control over the dynamics of such a symmetry-breaking impurity. Crucially, the topological protection of the boundary states would allow this control to be exerted over mesoscopic distances, as the excitations can propagate to the impurity with minimal loss.

In this chapter, we discuss the interaction between a TRS breaking impurity and the topological boundary states of a QSH model. Specifically, we couple a classical impurity spin $\mathbf{S}(t)$ to the boundary of a two-dimensional Kane–Mele (KM) model [20, 21]. The Kane–Mele model has originally been proposed for graphene [20] but turned out to be relevant for certain quantum-well systems [152, 153] too. It can also be understood to describe a class of graphene-like two-dimensional monolayer honeycomb materials that feature significant spin-orbit interaction, such as silicene and related systems [154, 155]. An interacting Kane–Mele model emerges as an effective low-energy theory in stacked 1T-TaSe2 bilayers [156]. The idea of probing TRS protected topological states by means of TRS breaking local perturbations [149] has been pursued in various studies. For example, by doping TRS invariant systems, such as Bi_2Te_3 , Bi_2Se_3 or Sb_2Te_3 , with magnetic transition-metal atoms [157, 158] or by depositing magnetic adatoms at the surface [159, 160]. Locally breaking TRS may also lead to rather exotic phenomena like an image magnetic monopole [161].

Another intriguing aspect is the interaction between two magnetic adatoms mediated by the helical edge states. For weak exchange couplings J between adatoms and substrate, standard RKKY theory [162] can be adapted to tight-binding or continuum models describing helical QSH boundary states. Near a classical magnetic impurity, the local (spin) density of states is suppressed at low energies [163]. When the chemical potential places the Fermi level such that the resulting Fermi wavelength exceeds the impurity separation, the RKKY interaction between two impurities becomes ferromagnetic. In general, the coupling is non-collinear, confined to the plane, and decays spatially as a power law [164]. In addition, there exists a Bloembergen–Rowland-type [165] bulk contribution that decays exponentially with distance [166]. A weak breaking of time-reversal symmetry gaps the Dirac cones and induces a strongly anisotropic RKKY coupling, which also decays exponentially and includes Dzyaloshinskii–Moriya terms alongside in-plane and out-of-plane Ising components [167]. The validity of RKKY theory is further limited by strong electron interactions and the Doniach competition between indirect exchange and Kondo screening in helical Luttinger liquids [168].

The *real-time dynamics* of magnetic impurities at the surfaces of topological insulators has been studied less extensively. Recent studies have employed time-dependent density-functional theory [169] and, for periodically driven impurities, Floquet theory [170]. Scattering theory has been applied to study the influence of individual TRS-breaking magnetic impurities on the transport properties of the helical edge states of Kane–Mele zigzag ribbons [171]. Apart from the study in Ref. [172], which investigates the long-time dynamics of a single classical spin exchange-coupled to the edge of a Su–Schrieffer–Heger model, the full microscopic real-time dynamics beyond the linear-response regime [173–175] remains largely unexplored.

In the following, we numerically study the real-time dynamics of a classical “read-out” spin $\mathbf{S}(t)$ as it interacts with the helical boundary states of the Kane–Mele model. The dynamical state of the read-out spin is affected by another impurity at a distant site on the same edge, which is used to inject a local spin excitation. The time-dependent transport of injected spin density through the helical edge states and its effect on the classical spin are analysed. The goal is to exploit the topological protection of the helical QSH edge state to achieve dynamical control over the classical spin state across mesoscopic distances. We demonstrate that this can be achieved by iterating the spin-injection and transport processes.

Our setup is in part motivated by the progress of experimental techniques, such as detecting states of magnetic adatoms [176] and measuring indirect magnetic RKKY interactions on a nanoscale [177]. The spin-momentum-locked transport explored here could, for instance, be experimentally accessed using scanning-tunnelling microscopy with multiple tips. Such setups would enable initiation, probing, and control of the dynamics, provided they can be spaced down to nanometer scales and are equipped with magnetic, spin-resolved capabilities. We aim to advance the understanding of the dynamical manipulation of local magnetic states via topological surface states. However, since current time-dependent STM techniques [178, 179] operate on the microsecond rather than picosecond timescale, the focus here is on the initial and final spin configurations.

In general, a numerically exact solution of the coupled set of equations of motion describing the classical read-out spin and the electronic system can only be achieved for systems of finite size. Here, we demonstrate that a ribbon-shaped geometry with only 8 sites perpendicular and about 100 sites parallel to the zigzag edge is sufficient to achieve long-time propagation of electronic excitations of up to $\sim 10^3$ inverse hoppings. This is made possible by applying Lindblad-type absorbing boundaries, as introduced in [172, 180], along three out of the four edges to suppress reflections.

The remainder of this chapter is organised as follows. In Sec. 8.1, we review the Kane–Mele model, concentrating on the fundamental band structure, its topological properties, and the bulk-boundary correspondence that gives rise to the helical edge states. Afterwards, in Sec. 8.2, we outline the equations of motions governing the dynamics of the quantum-classical hybrid system. Subsequently, we introduce a numerical model for simulating macroscopic edge dynamics in Sec. 8.3, and analyse the helical propagation of spin density injections in the presence and absence of a read-out spin in Secs. 8.4 and 8.5, respectively. Having discussed the electronic dynamics, we then turn to the dynamics of the classical read-out spin in Sec. 8.6. In Sec. 8.7, we present the numerical results demonstrating how elementary spin density injections effect the classical spin and how iterating these elementary injections can be used to implement a helical spin switch protocol.

Throughout this chapter, we closely follow our original presentation in [RQ2].

8.1 The Kane–Mele Model

The Kane–Mele model provides the first microscopic description of the QSH effect [20]. It extends the tight-binding model of graphene by including spin-orbit interactions. These open up a gap and enable a topologically non-trivial insulating phase: the quantum spin Hall (QSH) phase. The second-quantised Kane–Mele Hamiltonian reads

$$H_{\text{KM}} = -t_{\text{hop}} \sum_{\langle j,k \rangle} \sum_{\alpha} c_{j\alpha}^{\dagger} c_{k\alpha} + V \sum_{j,\alpha} \epsilon_j c_{j\alpha}^{\dagger} c_{j\alpha} + it_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle} \sum_{\alpha,\beta} \nu_{jk} \sigma_z^{\alpha\beta} c_{j\alpha}^{\dagger} c_{k\beta}, \quad (8.1)$$

where j and k label the L sites of a honeycomb lattice in the xy -plane, cf. Sec 7.2.1. The expressions $\langle j, k \rangle$ and $\langle\langle j, k \rangle\rangle$ indicate summation over pairs of NN and NNN sites, and $\alpha, \beta \in \{\uparrow, \downarrow\}$ denote the spin projection of the electrons along the z -axis. The first term of Eq. (8.1) is the generic tight-binding hopping of graphene. It is governed by the real NN hopping amplitude t_{hop} and preserves all spatial (lattice and z -reflection) symmetries and the SU(2) spin symmetry of the electrons. In the following, the lattice constant $a \equiv 1$ sets the length unit, while the NN hopping amplitude $t_{\text{hop}} \equiv 1$ sets the energy unit and, together with $\hbar \equiv 1$, also the time unit. The second term is a staggered sublattice potential. It is characterised by the real on-site potential strength V and the sign $\epsilon_j = \pm 1$, which is positive (negative) when j belongs to the A (B) sublattice. For $V \neq 0$, this term breaks the I_2 sublattice inversion symmetry ($I_2 \rightarrow 1$) and reduces the six-fold rotational symmetry C_6 to a three-fold rotational symmetry ($C_6 \rightarrow C_3$). The SU(2) spin symmetry and the z -reflection symmetry R_2 are left invariant. The third term describes intrinsic spin-orbit coupling of the electrons. It is determined by the real NNN spin-orbit hopping amplitude t_{SO} , the sign $\nu_{jk} = \pm 1$, which is positive (negative) for anticlockwise (clockwise) hopping $k \rightarrow j$ within a hexagon of the lattice, and the z -Pauli matrix σ_z , whose elements mediate the coupling between the spin projections. If $t_{\text{SO}} \neq 0$, this term reduces the electronic SU(2) spin symmetry to U(1) rotations around the z -axis ($\text{SU}(2) \rightarrow \text{U}(1)$). Furthermore, the symmetry under R_2 is preserved since $\sigma_z \rightarrow -\sigma_z$ is compensated by $\nu_{jk} \rightarrow -\nu_{jk}$, which happens because $z \rightarrow -z$ reverses the orientation of the xy -plane.¹

Note that the total hopping amplitude $it_{\text{SO}}\nu_{jk}\sigma_z^{\alpha\beta}$ of the intrinsic NNN spin-orbit hopping is always imaginary. In this respect, the Kane–Mele model is similar to the Haldane model, which adds complex hoppings to the tight-binding model of graphene to induce an intrinsic breaking of TRS and implement a quantum anomalous Hall phase [19]. However, unlike the Haldane model, the Kane–Mele model is *invariant* under TRS, fulfilling

$$\mathcal{T}H_{\text{KM}}\mathcal{T}^\dagger = H_{\text{KM}}, \quad (8.2)$$

where $\mathcal{T}$ denotes the TRS transformation of spin one-half fermions given in Eq. (3.15). A proof of Eq. (8.2) is provided in App. A.10. Recall that the TRS operator $\mathcal{T}$ of spin one-half fermions satisfies $\mathcal{T}^2 = -\mathbb{1}$, so that Eq. (8.2) causes the single-particle eigenstates of the Kane–Mele model to be Kramers degenerate. This has profound consequences for the topological properties of the Kane–Mele model, as we will see shortly. The Kane–Mele Hamiltonian H_{KM} can be diagonalised in $\mathbf{k}$ -space. Specifically, the Fourier transform $c_{j\alpha} = 1/\sqrt{L} \sum_{\mathbf{k}} e^{i\mathbf{k}\cdot\mathbf{R}_j} c_{\mathbf{k}\alpha}$ of the elementary field operators allows us to write H_{KM} as

$$H_{\text{KM}} = \sum_{\mathbf{k}} \phi^\dagger(\mathbf{k}) h_{\text{KM}}(\mathbf{k}) \phi(\mathbf{k}), \quad (8.3)$$

where we introduced the spinor $\phi(\mathbf{k}) = (a_{\mathbf{k}\uparrow} a_{\mathbf{k}\downarrow} b_{\mathbf{k}\uparrow} b_{\mathbf{k}\downarrow})^\top$ of annihilation operators $a_{\mathbf{k}\alpha}$ and $b_{\mathbf{k}\alpha}$ for Bloch states with quasi-momentum $\mathbf{k}$ and spin projection $\alpha \in \{\uparrow, \downarrow\}$ on the A and B sublattices, respectively. Again, the A and B sublattices form a two-component degree of freedom that behaves mathematically like a spin and called a sublattice pseudospin in the following. The spinors in Eq. (8.3) are then given in a $\sigma \otimes \tau$ tensor basis where σ describes the $\{\uparrow, \downarrow\}$ components of electron spin, while τ represents the $\{a, b\}$ components of sublattice pseudospin.

The Hermitian 4×4 Bloch matrix $h_{\text{KM}}(\mathbf{k})$ from Eq. (8.3) takes the form (for details see App. A.10)

$$h_{\text{KM}}(\mathbf{k}) = h_0(\mathbf{k}) \sigma_z \otimes \tau_z + \mathbb{1}_2 \otimes [\mathbf{h}(\mathbf{k})\boldsymbol{\tau}] = h_0(\mathbf{k}) \sigma_z \otimes \tau_z + \sum_{\mu=x,y,z} h_\mu(\mathbf{k}) \mathbb{1}_2^\sigma \otimes \tau_\mu, \quad (8.4)$$

which is given in terms of the electron-spin identity matrix $\mathbb{1}_2^\sigma$, the electron-spin z -Pauli matrix σ_z and the vector $\boldsymbol{\tau}$ of sublattice-pseudospin Pauli matrices. The functions $h_0(\mathbf{k})$ and $h_\mu(\mathbf{k})$ with $\mu = x, y, z$ are presented in Tab. 8.1. Note that the basis matrices in Eq. (8.4) transform as (for details see App. A.10)

$$\begin{aligned} \mathcal{T}(\sigma_z \otimes \tau_z) \mathcal{T}^\dagger &= -(\sigma_z \otimes \tau_z), & \mathcal{T}(\mathbb{1}_2 \otimes \tau_x) \mathcal{T}^\dagger &= (\mathbb{1}_2 \otimes \tau_x) \\ \mathcal{T}(\mathbb{1}_2 \otimes \tau_y) \mathcal{T}^\dagger &= -(\mathbb{1}_2 \otimes \tau_y), & \mathcal{T}(\mathbb{1}_2 \otimes \tau_z) \mathcal{T}^\dagger &= (\mathbb{1}_2 \otimes \tau_z) \end{aligned} \quad (8.5)$$

¹This can be seen by writing $\nu_{jk} = 2(\mathbf{n}_1 \times \mathbf{n}_2)_z / \sqrt{3}$ where $\mathbf{n}_1$ and $\mathbf{n}_2$ are the unit vectors along the two NN bonds the electron traverses to get from site k to site j . Clearly $(\mathbf{n}_1 \times \mathbf{n}_2)_z \rightarrow -(\mathbf{n}_1 \times \mathbf{n}_2)_z$, and thus $\nu_{jk} \rightarrow -\nu_{jk}$ under $z \rightarrow -z$.

$h_0(\mathbf{k})$	$-2t_{\text{SO}}(2 \sin x \cos 3y - \sin 2x)$	$h_x(\mathbf{k})$	$-t_{\text{hop}}(2 \cos x \cos y + \cos 2y)$
$h_y(\mathbf{k})$	$-t_{\text{hop}}(2 \cos x \sin y - \sin 2y)$	$h_z(\mathbf{k})$	V

Table 8.1: Non-zero coefficients of Eq. (8.3), with $x = \sqrt{3}a_0k_x/2$ and $y = a_0k_y/2$.

under TRS, so that the TRS invariance of the Kane–Mele Hamiltonian demands that

$$h_0(\mathbf{k}) = -h_0(-\mathbf{k}), \quad h_x(\mathbf{k}) = h_x(-\mathbf{k}), \quad h_y(\mathbf{k}) = -h_y(-\mathbf{k}), \quad h_z(\mathbf{k}) = h_z(-\mathbf{k}) \quad (8.6)$$

under the TRS-induced involutive transformation

$$\vartheta : \mathbb{T}_{\mathbf{k}}^2 \rightarrow \mathbb{T}_{\mathbf{k}}^2, \quad \mathbf{k} \mapsto -\mathbf{k} \quad (8.7)$$

on the two-dimensional Brillouin torus $\mathbb{T}_{\mathbf{k}}^2$. It can easily be verified that the coefficient functions from Tab. 8.1 readily satisfy these TRS constraints. The diagonalisation of the Bloch matrix $h_{\text{KM}}(\mathbf{k})$ in Eq. (8.4) yields four energy bands

$$E_{\alpha}^{\pm}(\mathbf{k}) = \pm \sqrt{h_x(\mathbf{k})^2 + h_y(\mathbf{k})^2 + (h_z(\mathbf{k}) + \eta_{\alpha}h_0(\mathbf{k}))^2}, \quad (8.8)$$

where $\eta_{\alpha} = \pm 1$ is positive for spin-up ($\alpha = \uparrow$) and negative for spin-down ($\alpha = \downarrow$). For each spin projection $\alpha \in \{\uparrow, \downarrow\}$ we therefore get two bands, a conduction band $E_{\alpha}^{+}(\mathbf{k})$ and a valence band $E_{\alpha}^{-}(\mathbf{k})$, which are symmetric around zero. As a result, the energy gap of the system is defined as

$$\Delta E := \min_{\alpha, \mathbf{k}} (E_{\alpha}^{+}(\mathbf{k}) - E_{\alpha}^{-}(\mathbf{k})) = 2 \cdot \min_{\alpha, \mathbf{k}} (E_{\alpha}^{+}(\mathbf{k})), \quad (8.9)$$

and we find that the minimum is attained at the Dirac points

$$\mathbf{K}^{\pm} = \pm \frac{4\pi}{3\sqrt{3}a_0} \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad (8.10)$$

where $h_x(\mathbf{K}^{\pm}) = h_y(\mathbf{K}^{\pm}) = 0$ and $h_0(\mathbf{K}^{\pm}) = \mp 3\sqrt{3}t_{\text{SO}}$, so that

$$\Delta E = 2 \cdot \min_{\mathbf{K}^{\pm}} |V \mp 3\sqrt{3}t_{\text{SO}}|. \quad (8.11)$$

Note that the bulk gap ΔE closes along the nodal line $|V| = 3\sqrt{3}|t_{\text{SO}}|$ in the two-dimensional parameter space spanned by t_{SO} and V . This nodal line divides the parameter space into two separate regions with finite gaps $\Delta E > 0$: one dominated by the onsite potential ($|V| > 3\sqrt{3}|t_{\text{SO}}|$) and one dominated by the intrinsic spin-orbit coupling ($|V| < 3\sqrt{3}|t_{\text{SO}}|$). It turns out that these two regions correspond directly to the topologically trivial and non-trivial phases of the Kane–Mele model. For this reason, the onsite potential V is often used to tune between the topologically distinct phases in practice. Here, we typically choose t_{SO} and ΔE freely and then define

$$V = 3\sqrt{3}t_{\text{SO}} \pm \Delta E/2 \quad (8.12)$$

to generate a topologically trivial (positive sign) or a topologically non-trivial (negative sign) band structure with the same chosen bulk band gap of $\Delta E > 0$.

At half-filling, i.e. for every chemical potential μ with $|\mu| < \Delta E/2$, the ground state $|\text{GS}\rangle$ of the Kane–Mele model is a Slater determinant

$$|\text{GS}\rangle = \prod_{\substack{\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^2 \\ \alpha = \uparrow, \downarrow}} d_{\mathbf{k}\alpha}^{-\dagger} |0\rangle = \bigwedge_{\substack{\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^2 \\ \alpha = \uparrow, \downarrow}} |u_{\alpha}^{-}(\mathbf{k})\rangle \quad (8.13)$$

of all valence Bloch states $|u_{\alpha}^{-}(\mathbf{k})\rangle \equiv d_{\mathbf{k}\alpha}^{-\dagger} |0\rangle \in \mathcal{H} \subset \mathcal{F}$. Here, $\mathcal{H} \subset \mathcal{F}$ indicates the natural inclusion of the single-particle Hilbert space $\mathcal{H}$ into the many-particle Fock space $\mathcal{F}$ and $|0\rangle$ denotes the vacuum state of $\mathcal{F}$ defined by $c_{\mathbf{k}\alpha} |0\rangle = 0$ for all $\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^2$ and $\alpha \in \{\uparrow, \downarrow\}$. The topology of the valence Bloch states therefore determines the topological properties of the many-body ground state.

8.1.1 Topology of the Kane–Mele Model

The family $\{\mathcal{H}(\mathbf{k})\}_{\mathbf{k} \in \mathbb{T}_k^2}$ of Bloch spaces

$$\mathcal{H}(\mathbf{k}) := \text{span}(|u_{\uparrow}^+(\mathbf{k})\rangle, |u_{\downarrow}^+(\mathbf{k})\rangle, |u_{\uparrow}^-(\mathbf{k})\rangle, |u_{\downarrow}^-(\mathbf{k})\rangle) \quad (8.14)$$

defines a rank-four Bloch bundle $\mathcal{B}_{\text{KM}} \xrightarrow{\pi} \mathbb{T}_k^2$ over the two-dimensional Brillouin torus $\mathbb{T}_k^2$. For $\Delta E > 0$, this Bloch bundle can be split as

$$\mathcal{B}_{\text{KM}} = \mathcal{B}_{\text{KM}}^- \oplus \mathcal{B}_{\text{KM}}^+, \quad (8.15)$$

where $\mathcal{B}_{\text{KM}}^{\pm} \xrightarrow{\pi^{\pm}} \mathbb{T}_k^2$ are the two rank-two valence and conduction subbundles of $\mathcal{B}_{\text{KM}}$ that are determined by the families $\{\mathcal{H}^{\pm}(\mathbf{k})\}_{\mathbf{k} \in \mathbb{T}_k^2}$ of valence and conduction subspaces

$$\mathcal{H}^{\pm}(\mathbf{k}) := \text{span}(|u_{\uparrow}^{\pm}(\mathbf{k})\rangle, |u_{\downarrow}^{\pm}(\mathbf{k})\rangle). \quad (8.16)$$

This is again similar to the spinless Haldane model, where the rank-two Bloch bundle $\mathcal{B}_{\text{H}}$ splits into two rank-one valence and conduction subbundles $\mathcal{B}_{\text{H}}^{\pm}$, whose topology is ultimately characterised by the first Chern number $C_1^{(k)}$, cf. Eq. (7.33). Based on these parallels, it is natural to wonder whether an analogous topological classification is possible for the Kane–Mele Bloch bundle $\mathcal{B}_{\text{KM}}$ and its valence and conduction subbundles $\mathcal{B}_{\text{KM}}^{\pm}$. It turns out that this is not the case. For one thing, the indispensable inclusion of spin in the Kane–Mele model promotes $\mathcal{B}_{\text{KM}}^{\pm}$ to complex vector bundles of rank two. In order to repeat the Chern classification of $\mathcal{B}_{\text{H}}^{\pm}$ for $\mathcal{B}_{\text{KM}}^{\pm}$ we would therefore have to further split the rank-two bundles $\mathcal{B}_{\text{KM}}^{\pm}$ into two rank-one bundles $\mathcal{B}_{\text{KM},\uparrow}^{\pm}$ and $\mathcal{B}_{\text{KM},\downarrow}^{\pm}$ of spin-polarised valence and conduction bands. However, this is not generally possible because Kramers degeneracy demands that

$$E_{\uparrow}^{\pm}(\mathbf{k}) = E_{\downarrow}^{\pm}(-\mathbf{k}), \quad (8.17)$$

i.e. that the energy $E_{\uparrow}^{\pm}(\mathbf{k})$ of each Bloch state $|u_{\uparrow}^{\pm}(\mathbf{k})\rangle$ is the same as the energy $E_{\downarrow}^{\pm}(-\mathbf{k})$ of its TRS partner $\mathcal{T}|u_{\uparrow}^{\pm}(\mathbf{k})\rangle = |u_{\downarrow}^{\pm}(-\mathbf{k})\rangle$. A direct consequence of this is that the spin-up and the spin-down valence and conduction bands are *glued together*,

$$E_{\uparrow}^{\pm}(\boldsymbol{\kappa}_j) = E_{\downarrow}^{\pm}(\boldsymbol{\kappa}_j), \quad (8.18)$$

at the four TR invariant (quasi-)momenta (TRIM)

$$(\boldsymbol{\kappa}_1, \boldsymbol{\kappa}_2, \boldsymbol{\kappa}_3, \boldsymbol{\kappa}_4) = (\boldsymbol{\Gamma}, \boldsymbol{M}, \boldsymbol{M}', \boldsymbol{M}'') \equiv \mathcal{K}, \quad (8.19)$$

that are mapped to their own quasi-momentum equivalence class, $\vartheta(\boldsymbol{\kappa}_j) = -\boldsymbol{\kappa}_j \simeq \boldsymbol{\kappa}_j$, under the TRS-induced involution ϑ from Eq. (8.7). Of course, one could simply attempt to base the bundle classification on the second Chern number of the rank-two valence and conduction bundle instead. However, this does not work either. The reason is that TRS with $\mathcal{T}^2 = -1$ twists the Bloch bundle in a way that makes its Chern numbers vanish [14]. To illustrate this, note that the TRS operator $\mathcal{T}$ permutes the fibres of the valence and conduction subbundles as

$$\mathcal{T}|u_{\uparrow}^{\pm}(\mathbf{k})\rangle = |u_{\downarrow}^{\pm}(-\mathbf{k})\rangle \quad \text{and} \quad \mathcal{T}|u_{\downarrow}^{\pm}(\mathbf{k})\rangle = -|u_{\uparrow}^{\pm}(-\mathbf{k})\rangle, \quad (8.20)$$

intertwining the spin-up and spin-down valence and conduction bands across the two-dimensional Brillouin torus. In the mathematical literature, TRS with $\mathcal{T}^2 = -1$ is said to impose a quaternionic structure

$$I = i1, \quad J = \mathcal{T}, \quad K = i\mathcal{T} \quad \text{with} \quad I^2 = J^2 = K^2 = IJK = -1 \quad (8.21)$$

on the Bloch bundle [181]. The complex rank-two valence and conduction subbundles $\mathcal{B}_{\text{KM}}^{\pm}$ can then be understood as quaternionic rank-one bundles. These are classified within the framework of quaternionic K -theory, or KQ -theory, and it is found that quaternionic line-bundles have a $\mathbb{Z}_2$ classification on the two-dimensional (Brillouin) torus. This shows that the topology of the Kane–Mele Bloch bundle is fundamentally different from that of the Haldane Bloch bundle, whose Chern number generates a $\mathbb{Z}$ classification on the two-dimensional (Brillouin) torus.

There are various formulas for the $\mathbb{Z}_2$ invariant ν characterising the topology of the Kane–Mele valence subbundle $\mathcal{B}_{\text{KM}}^-$. In their original publication [20, 21], Kane and Mele propose a formula that can be written as

$$(-1)^\nu = \prod_{\kappa_j \in \mathcal{K}} \text{sign}[\text{Pf}(\omega(\kappa_j))]. \quad (8.22)$$

Here, $\kappa_j \in \mathcal{K}$ are the four TRIM from Eq. (8.19), $\text{Pf}(\cdot)$ denotes the Pfaffian, and $\omega(\kappa_j)$ is a skew-symmetric 2×2 TRS scattering matrix with elements

$$\omega_{rs}(\kappa_j) = \langle u_r(\kappa_j) | \mathcal{T} | u_s(\kappa_j) \rangle, \quad (8.23)$$

where the indices $r, s = 1, 2$ label the two Bloch states $|u_1(\kappa)\rangle = |u_\uparrow^-(\kappa)\rangle$ and $|u_2(\kappa)\rangle = |u_\downarrow^-(\kappa)\rangle$ of the valence Kramers pair [181]. A more detailed construction of Eq. (8.22) is given in App. A.10. Despite the fact that Kramers degeneracy prevents a trivial splitting

$$\mathcal{B}_{\text{KM}}^\pm = \mathcal{B}_{\text{KM},\uparrow}^\pm \oplus \mathcal{B}_{\text{KM},\downarrow}^\pm \quad (8.24)$$

of the rank-two bundles $\mathcal{B}_{\text{KM}}^\pm$ into spin-polarised rank-one subbundles $\mathcal{B}_{\text{KM},\uparrow}^\pm$ and $\mathcal{B}_{\text{KM},\downarrow}^\pm$ in general, such a decomposition can still become possible in certain cases. Specifically, if there are no interactions between the two spin projections, i.e. if the only spin interactions present in the model are spin-diagonal, like the intrinsic spin-orbit term in Eq. (8.1), the spin-up and spin-down components of the valence and conduction bundle decouple and can be treated as independent Haldane-like rank-one bundles after all. In this case, their Chern numbers are defined as in Eq. (7.33) and may be non-zero because $\mathcal{B}_{\text{KM},\uparrow}^\pm$ and $\mathcal{B}_{\text{KM},\downarrow}^\pm$ are not *individually* invariant under TRS. However, even if the Chern numbers of the spin-polarised rank-one subbundles turn out non-zero, their values are not independent of each other: the Whitney sum formula Eq. (2.209) of the Chern classes tells us that

$$c_1(\mathcal{F}_{\text{KM}}^\pm) = c_1(\mathcal{F}_{\text{KM},\uparrow}^\pm \oplus \mathcal{F}_{\text{KM},\downarrow}^\pm) = c_1(\mathcal{F}_{\text{KM},\uparrow}^\pm) + c_1(\mathcal{F}_{\text{KM},\downarrow}^\pm), \quad (8.25)$$

and, hence, that

$$C_{1,\pm}^{(k)} = C_{1,(\pm,\uparrow)}^{(k)} + C_{1,(\pm,\downarrow)}^{(k)}. \quad (8.26)$$

Since the TRS invariance of the full rank-two valence and conduction subbundle $\mathcal{B}_{\text{KM}}^\pm$ enforces

$$C_{1,\pm}^{(k)} = 0, \quad (8.27)$$

we find that $C_{1,(\pm,\uparrow)}^{(k)}$ and $C_{1,(\pm,\downarrow)}^{(k)}$ must always cancel out, even if they are non-zero individually. This constraint motivates the definition

$$\nu := \frac{1}{2}(C_{1,(\pm,\uparrow)}^{(k)} - C_{1,(\pm,\downarrow)}^{(k)}) \mod 2 \quad (8.28)$$

of a topological $\mathbb{Z}_2$ invariant that distinguishes between even and odd Chern numbers $C_{1,(\pm,\uparrow)}^{(k)}$ and $C_{1,(\pm,\downarrow)}^{(k)}$. In the absence of interactions between the two spin projections, Eq. (8.28) provides another formula for the $\mathbb{Z}_2$ invariant of the Kane–Mele model.

For the sake of completeness, we note that there exists a generic TRS preserving spin-orbit coupling term that *does* mix different spin projections. Rashba spin-orbit coupling arises in the presence of electric fields, which break structural inversion symmetry. Such fields usually appear at surfaces or interfaces, and are especially relevant in monolayer materials like graphene. In these systems, a perpendicular electric field couples to the electron spin through its in-plane motion, leading to a spin-orbit term of the form

$$H_{\text{RSO}} = t_{\text{RSO}} (\mathbf{p} \times \boldsymbol{\sigma})_z = t_{\text{RSO}} (p_x \sigma_y - p_y \sigma_x), \quad (8.29)$$

where t_{RSO} is the Rashba spin-orbit coupling strength, $\mathbf{p}$ represents the electron momentum, and $\boldsymbol{\sigma}$ denotes the vector of Pauli matrices describing electron spin. The non-trivial ($\nu = 1$) QSH phase of the Kane–Mele model remains stable under the inclusion of weak Rashba spin-orbit coupling [20, 21]. However, once $t_{\text{RSO}} > 0$, the $\mathbb{Z}_2$ invariant can no longer be computed using Eq. (8.28).

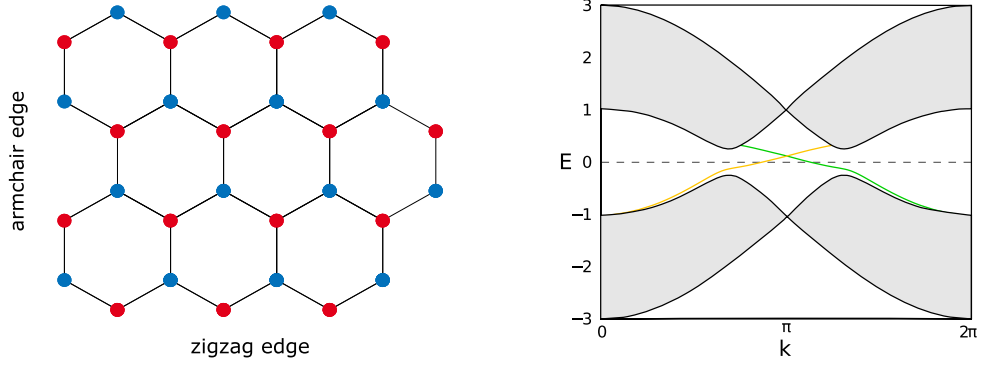


Figure 8.1: *Left:* Zigzag and armchair edges on a finite honeycomb tile. *Right:* Bandstructure of the topological Kane–Mele model, Eq. (8.1), projected onto one edge of a zigzag ribbon. The projected bulk bands, shown in grey, are obtained from a ribbon geometry calculation with a width of 20 unit cells in the perpendicular direction. Dispersions $\varepsilon_{\uparrow}(k)$ (yellow) and $\varepsilon_{\downarrow}(k)$ (green) of the two helical edge states outside the bulk continuum; edge states localised at the opposite edge are not shown. The NN hopping amplitude $t_{\text{hop}} = 1$ sets the energy scale. The remaining parameters were set to $t_{\text{SO}} = 0.05$, $\Delta E = 0.3$ and $V = 3\sqrt{3}t_{\text{SO}} - \Delta E/2$ with the negative sign. Reproduced with minor modifications from Ref. [RQ2].

8.1.2 Bulk-Boundary Correspondence in Zigzag Ribbons

Consider a Kane–Mele model on an infinite honeycomb lattice. In principle, we can introduce an edge to the system by cutting it along any direction. However, due to the symmetry of the lattice, certain edge directions are more natural than others. In the honeycomb lattice, the two canonical edge orientations are known as zigzag and armchair edges, cf. Fig. 8.1.

A honeycomb model with a single edge is typically referred to as a honeycomb *half-space model*. The edge breaks translation invariance in the direction perpendicular to it, but preserves translation invariance along the parallel direction. As a result, there is a one-dimensional parallel crystal momentum $k \in \mathbb{T}_k^1$ and we get an infinite number of energy bands dispersing over k . Half-space models capture the essential physics of a single boundary in an otherwise infinite system, so they provide a good description of the edges of macroscopic systems. A honeycomb half-space model that terminates on a zigzag (armchair) edge is called a zigzag (armchair) half-space model, respectively.

If we introduce two parallel edges to a system, it takes on a so-called *ribbon geometry*: infinite and periodic along the edge direction, but finite and open in the perpendicular direction. As before, translational invariance parallel to the edge gives rise to a one-dimensional parallel crystal momentum $k \in \mathbb{T}_k^1$, while the finite width perpendicular to the edges determines the number of sites in the unit cell and thus the local Hilbert space dimension d . Combined, we get d energy bands $E_j(k)$ with $j = 1, \dots, d$ that disperse over the parallel momentum. Ribbons with zigzag (armchair) edges are called zigzag (armchair) ribbons, respectively.

In zigzag ribbons, the parallel momentum always intersects the K and K' points of the projected bulk Brillouin zone. These are the points where pristine graphene has its band closures. As a result, zigzag ribbons are always metallic, whereas armchair ribbons can be metallic or semiconducting depending on the width of the ribbon [182]. This poses a practical complication for our numerical treatment. In the upcoming sections, we are going to be interested in ribbon and half-space Kane–Mele models. These are obtained from finite Kane–Mele tiles through partial Fourier transform and absorbing boundary conditions, respectively. The width of the finite tile in the direction perpendicular to the selected “physical” edge then sets the width of the resulting geometry. The fact that the presence or absence of zero modes in armchair ribbons depends on the ribbon width shows that the physics at armchair edges is quite sensitive to the geometric details along the direction perpendicular to the edge. Thus, even small changes in the tile geometry could have a big impact on the edge physics of the corresponding armchair ribbon or armchair-terminated half-space model. This would force us to distinguish between cases. To avoid elaborate case analyses and get generic results that are qualitatively independent of tile width, we focus on *zigzag ribbons* and *zigzag-terminated half-space models* in the following.

If we consider a Kane–Mele model with edges, the physical properties of the edges are, in part, determined by the bulk–boundary correspondence of the Kane–Mele model. In particular, a non-trivial bulk topology ($\nu = 1$) gives rise to helical edge states whenever the ribbon interfaces with a topologically trivial ($\nu = 0$) region, such as the vacuum. However, not all edge states are topological. In order to distinguish between accidental and topological edge states, we note that TRS forces any collection of (topological or accidental) edge states to form (at least) twofold degenerate Kramers pairs at each TRIM. Away from the TRIM, this degeneracy is generally lifted by the intrinsic spin-orbit coupling t_{SO} . Now there are two ways in which the edge states can connect Kramers pairs at distinct TRIM κ_1 and κ_2 : (a) each Kramers pair at κ_1 connects back to itself at κ_2 , or (b) the partners switch and connect to one another.² In scenario (a), the edge states cross the Fermi energy E_F an even number of times. In this case, they can be continuously pushed out of the bulk gap without closing it, and the system describes a trivial insulator. In scenario (b), the edge states cross the Fermi energy E_F an odd number of times. In this case, they are topologically protected and cannot be removed without closing the bulk gap. The topological phases of the Kane–Mele model on a zigzag ribbon³ are therefore characterised by the number of intersections N_F between its edge states and the Fermi energy: if N_F is even, the model is trivial, if N_F is odd, it is non-trivial. The bulk–boundary correspondence of the Kane–Mele model can therefore be expressed as

$$\Delta\nu = N_F \mod 2, \quad (8.30)$$

where $\Delta\nu = |\nu_{\text{ribbon}} - \nu_{\text{ambient}}|$ is the difference between the topological Kane–Mele $\mathbb{Z}_2$ invariants of the ribbon and the adjacent region [183]. The right picture in Fig. 8.1 shows an example of this. It displays the band structure of a topological ($\nu = 1$) Kane–Mele ribbon with zigzag edges in vacuum ($\nu = 0$). The projected bulk bands are indicated in grey and the spin-polarised edge states on one of the two zigzag edges are shown in yellow (spin-down) and green (spin-up) respectively. The two edge states are connected to each other at $k = 0$ and $k = \pi$, forming a Kramers pair that intersects the Fermi energy E_F an odd number of $\Delta\nu = N_F = 1$ times. Another feature that can be determined from Fig. 8.1 is the helicity, or spin-momentum locking, of the topological edge modes. To this end, we consider the Fermi velocity

$$v_F := \left. \frac{\partial \varepsilon_\alpha(k)}{\partial k} \right|_{k=k_F} \quad (8.31)$$

of the spin polarised edge states at the Fermi energy. It shows that the slope of the edge mode’s energy dispersion $\varepsilon_\alpha(k)$ determines their direction of propagation: under the convention that a positive slope corresponds to right-moving excitations, a negative slope indicates left-moving excitations. Figure 8.1 demonstrates that there are only right moving (positive slope) spin-down and left moving (negative slope) spin-up states at the selected edge⁴ of the ribbon. For this reason, we call the topological edge states of the QSH *spin-momentum locked* or *helical*. Note that we can even read off the Fermi velocities $v_F \approx \pm 0.285$ from Fig. 8.1, which is in good agreement with $v_F \approx \pm 0.286$ as obtained from an analytical expression given in [184].

The helicity of the QSHE boundary modes distinguishes them from the chiral boundary modes of the QHE. This is illustrated in Fig. 8.2. While the topological $\mathbb{Z}$ invariant characterising the integer QHE gives rise to a chiral edge mode that carries charge in a fixed direction around the boundary, the topological $\mathbb{Z}_2$ invariant of the QSHE produces no such a charge motion. Instead, it induces two counterpropagating edge currents; one for each spin species. While the *charge* transport of these currents cancels out exactly, they still generate non-trivial counterpropagating *spin-currents*. Thus, the helical edge states of the Kane–Mele model transport spin rather than charge.

²The notion of switching partners at TRIM also plays a role in the derivation of the Kane–Mele $\mathbb{Z}_2$ invariant presented in App. A.10.

³The same argument applies to the aforementioned setup of a half-space lattice, i.e. a system with just a single edge, given the edge-termination is compatible with the topological lattice model [183].

⁴The situation is reversed on the other edge of the ribbon.

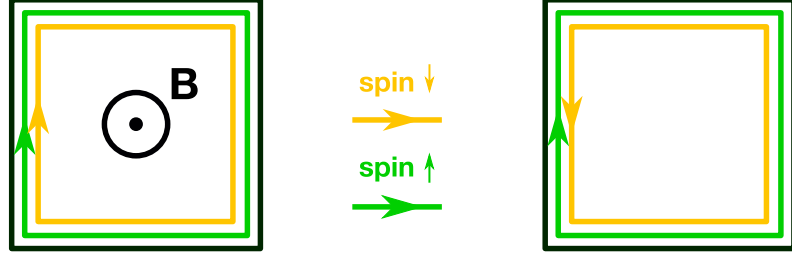


Figure 8.2: Qualitative comparison between the integer quantum Hall effect (left) and the quantum spin Hall effect (right) edge modes. The spin-up (spin-down) currents are sketched in green (yellow) colour. The quantum Hall setup features a TRS-breaking magnetic field $\mathbf{B}$ (black $\odot$). The spin currents in the quantum Hall (quantum spin Hall) system round the sample in the same (opposite) directions, giving rise to charge (spin) currents.

8.2 Real-Time Dynamics in the Kane–Mele s-d-Model

In order to study the dynamical interaction between the helical QSH edge states and a locally TRS breaking impurity, we first construct a suitable quantum-classical hybrid system based on the Kane–Mele model. Specifically, we consider the Kane–Mele model on a finite honeycomb tile with armchair and zigzag edges and add a Kondo-like perturbation term

$$H_R = J \mathbf{S}_R \mathbf{s}_R, \quad (8.32)$$

which implements a local exchange coupling between a classical (“read-out”) impurity spin $\mathbf{S}_R$ and the local magnetic moment $\mathbf{s}_R$ of the electron system at some site R on one of the zigzag edges. The latter is given in terms of its components

$$s_{R\mu} = \frac{1}{2} \sum_{\alpha,\beta} c_{R\alpha}^\dagger \sigma_\mu^{\alpha\beta} c_{R\beta}, \quad (8.33)$$

where $\mu = x, y, z$. The exchange coupling strength is denoted J and σ are the Pauli matrices describing electron spin. Importantly, the Kondo-like perturbation from Eq. (8.32) breaks TRS, since

$$\begin{aligned} \mathcal{T} s_{R\mu} \mathcal{T}^\dagger &\stackrel{(\diamond)}{=} \frac{1}{2} \sum_{\alpha,\beta} \mathcal{T} c_{R\alpha}^\dagger \mathcal{T}^\dagger \sigma_\mu^{*\alpha\beta} \mathcal{T} c_{R\beta} \mathcal{T}^\dagger \\ &\stackrel{(\star)}{=} \frac{1}{2} \sum_{\substack{\alpha,\beta \\ \gamma,\eta}} (-1)^{\delta_{\mu y}} (i \sigma_y^{\alpha\gamma}) c_{R\gamma}^\dagger \sigma_\mu^{\alpha\beta} (i \sigma_y^{\beta\eta}) c_{R\eta} \\ &\stackrel{(*)}{=} \frac{1}{2} \sum_{\substack{\alpha,\beta \\ \gamma,\eta}} (-1)^{\delta_{\mu y}} c_{R\gamma}^\dagger \sigma_y^{\gamma\alpha} \sigma_\mu^{\alpha\beta} \sigma_y^{\beta\eta} c_{R\eta} \\ &\stackrel{(\Delta)}{=} -\frac{1}{2} \sum_{\gamma,\eta} (-1)^{\delta_{\mu y}} (-1)^{\delta_{\mu y}} c_{R\gamma}^\dagger \sigma_\mu^{\gamma\eta} c_{R\eta} \\ &= -\frac{1}{2} \sum_{\gamma,\eta} c_{R\gamma}^\dagger \sigma_\mu^{\gamma\eta} c_{R\eta} \\ &= -s_{R\mu}, \end{aligned} \quad (8.34)$$

so that

$$\mathcal{T} H_R \mathcal{T}^\dagger = J \sum_{\mu} S_{R\mu} \mathcal{T} s_{R\mu} \mathcal{T}^\dagger = -J \sum_{\mu} S_{R\mu} s_{R\mu} = -H_R. \quad (8.35)$$

In Eq. (8.34), we first used the antiunitarity of $\mathcal{T}$ to rewrite $\sigma_\mu^{\alpha\beta} = \mathcal{T}^\dagger \mathcal{T} \sigma_\mu^{\alpha\beta} = \mathcal{T}^\dagger \sigma_\mu^{*\alpha\beta} \mathcal{T}$ in $(\diamond)$. Then, we plugged in the TRS transformation of spin one-half fermions from Eq. (3.15) and introduced the shorthand notation $\sigma_\mu^* = (-1)^{\delta_{\mu y}} \sigma_\mu$ for the complex conjugate of the Pauli matrices σ_μ with $\mu = x, y, z$

in $(\star)$. Next, we cancelled $i^2 = -1$ with the sign that the y -Pauli matrix acquires under transposition, i.e. $\sigma_y^{\alpha\gamma} = -\sigma_y^{\gamma\alpha}$ in $(\star)$. Finally, we applied the transformation behaviour $\sigma_y^{\gamma\alpha} \sigma_\mu^{\alpha\beta} \sigma_y^{\beta\eta} = -(-1)^{\delta_{\mu y}} \sigma_\mu^{\gamma\eta}$ of the Pauli matrices σ_μ with $\mu = x, y, z$ under conjugation by σ_y in (Δ) . It is worth mentioning that even though Eq. (8.32) bears a formal resemblance of a Kondo impurity term, it describes a conceptually different perturbation. In particular, the fact that $\mathbf{S}_R$ is a classical spin means that there is no Kondo screening and thus no Kondo effect [185]. Moreover, a conventional Kondo coupling to a quantum-spin one-half $\mathbf{s}^q$ would not break TRS since $\mathcal{T} \mathbf{s}^q \mathcal{T}^\dagger = -\mathbf{s}^q$, similar to Eq. (8.34). This makes the classical impurity spin $\mathbf{S}_R$ the most natural choice for the TRS breaking magnetic impurity. Physically, classical spins provide a reasonable model for magnetic adatoms, i.e. adatoms with a well-defined spin moment that remains stable over all other relevant timescales of the system.

Now, the classical “read-out” spin $\mathbf{S}_R$ is susceptible to spin-density excitations propagating along the edge in the helical edge states of the topological Kane–Mele model. In order to inject a local spin excitation into the edge channels, we make use of another local TRS-breaking perturbation of the form Eq. (8.32),

$$H_I = -\mathbf{B}_I \mathbf{s}_I, \quad (8.36)$$

where the local magnetic field $\mathbf{B}_I$ is thought of as an externally aligned⁵ magnetic adatom. The injection site I is chosen to be far away from the read out site R to avoid direct interaction. In practice, the injection is facilitated by abruptly switching $\mathbf{B}_I$ on and off.

The total Hamiltonian of the hybrid system, consisting of the Kane–Mele ribbon, the exchange-coupled “read-out” spin $\mathbf{S}_R$ at edge site R and the local magnetic “injection” field B_I at edge site I , then reads

$$H = H_{\text{KM}} + H_R + H_I, \quad (8.37)$$

with the individual terms H_{KM} , H_R and H_I as specified in Eq. (8.1), Eq. (8.32) and Eq. (8.36), respectively. Note that the total Hamiltonian H represents a two-impurity s - d -type model [186].

Our goal is to fully determine the microscopic real-time dynamics of the hybrid system described by Eq. (8.37). To achieve this, we introduce the one-particle reduced density matrix $\rho(t)$ with elements

$$\rho_{(j\alpha)(k\beta)}(t) = \langle \Psi(t) | c_{k\beta}^\dagger c_{j\alpha} | \Psi(t) \rangle, \quad (8.38)$$

where $|\Psi(t)\rangle$ is the N -particle quantum state of the electron system at time t . This immediately implies $\text{tr} \rho = N$. We consider a half-filled system where $N = L$, with the number of lattice sites L . For quantum-classical hybrid systems [187, 188] of the form in Eq. (8.37), there exists a closed system of equations of motion [189], consisting of a Landau-Lifschitz-type equation,

$$\frac{d\mathbf{S}_R(t)}{dt} = J \langle \mathbf{s}_R(t) \rangle \times \mathbf{S}_R(t), \quad (8.39)$$

with $\langle \mathbf{s}_R(t) \rangle = \frac{1}{2} \sum_{\alpha\beta} \rho_{(R\alpha)(R\beta)}(t) \boldsymbol{\sigma}^{\beta\alpha}$ for the read-out spin, and a von Neumann equation,

$$i \frac{d\rho(t)}{dt} = [T^{\text{eff}}(t), \rho(t)], \quad (8.40)$$

for the density matrix. Here, $T^{\text{eff}}(t)$ is the effective hopping matrix with elements given by

$$T_{(j\alpha)(k\beta)}^{\text{eff}}(t) = T_{(j\alpha)(k\beta)} + \delta_{jR} \delta_{kR} J \frac{1}{2} (\mathbf{S}_R(t) \boldsymbol{\sigma})^{\alpha\beta} - \delta_{jI} \delta_{kI} \frac{1}{2} (\mathbf{B}_I(t) \boldsymbol{\sigma})^{\alpha\beta}, \quad (8.41)$$

where

$$T_{(j\alpha)(k\beta)} = -t_{\text{hop}} \delta_{\langle j,k \rangle} \delta_{\alpha\beta} + V \epsilon_j \delta_{jk} \delta_{\alpha\beta} + i t_{\text{SO}} \delta_{\langle j,k \rangle} \nu_{jk} \sigma_z^{\alpha\beta} \quad (8.42)$$

are the elements of the hopping matrix T of the pristine Kane–Mele model from Eq. (8.1). Using standard numerical methods for systems of ordinary differential equations, we can solve these equations of motion for finite systems with up to $L \approx 10^3$ sites [189, 190].

⁵This may, for instance, be achieved by means of a spin-polarized STM tip.

In a finite system with open boundary conditions, an edge excitation injected at a distance d from the nearest corner will generally propagate along the edge with Fermi velocity v_F and reflect off the corner after a time $t_{\text{ref}} \approx d/|v_F|$. The backscattered excitations then travel back along the edge, potentially interfering with the dynamics of an edge impurity. In order to avoid such finite-size induced interference, we either have to limit ourselves to very short time scales, consider very large systems, or find a way to suppress boundary reflections directly. Here, we follow the latter approach and adopt absorbing boundary conditions of the type introduced in Ref. [180]. These have been shown to effectively eliminate dynamical finite-size effects, enabling the study of real-time dynamics on long time scales even for systems of moderate size. Following Ref. [180], we thus extend Eq. (8.40) as

$$\frac{d\rho(t)}{dt} = -i [T^{\text{eff}}(t), \rho(t)] - \{\gamma, \rho(t) - \rho(0)\}, \quad (8.43)$$

where $\{\cdot, \cdot\}$ denotes the anticommutator and where γ is a non-negative diagonal matrix that defines the interaction between the system and an external dissipative bath. Here, the elements of γ are chosen as

$$\gamma_{(j\alpha)(k\beta)} = \delta_{jk} \delta_{\alpha\beta} \gamma_{j\alpha}, \quad (8.44)$$

and the vector of non-negative real numbers $\gamma_{i\alpha}$ fully determines the coupling strengths between the individual single-particle states and the bath environment. In order to effectively simulate a half-space model with a single zigzag edge, we choose $\gamma_{i\alpha}$ to be non-zero only for single-particle states associated with sites in a thin shell located at three out of the four edges of the finite-size honeycomb tile, namely the two armchair edges and one of the zigzag edges. A systematic comparison of the Lindblad dynamics with short-time open-boundary dynamics showed that a uniform Lindblad shell of unit thickness and spin-independent coupling ($\gamma_{j\alpha} = \gamma_j \equiv \gamma_0$) is sufficient to realise absorption without affecting the dynamics in the impurity region. Consequently, the Lindblad bath is governed by a single scalar parameter γ_0 , which is set $\gamma_0 = 0.2$ throughout this study, cf. Refs. [172, 180]. A sketch of this Lindblad configuration can be found in Fig 8.3, where the Lindblad sites are highlighted in grey colour. With this selective bath coupling, we effectively simulate the dynamics of a zigzag-terminated half-space model using the numerically much more accessible architecture of a finite Kane–Mele tile.

It is worth mentioning that Eq. (8.44) can be derived from a general Lindblad master equation [191, 192] by restricting the theory to non-interacting electrons. This results in an equation of motion involving only the one-particle reduced density matrix, rather than the full many-body statistical operator. Furthermore, Eq. (8.44) features an additional term proportional to the initial state $\rho(0)$. This term is not present in the standard Lindblad formalism, but required to prevent unwanted excitations *generated* by the Lindblad bath. This was demonstrated in Refs. [172, 180] for one-dimensional tight-binding systems with classical-spin impurities. Here, we adapt the method to the two-dimensional geometry of the Kane–Mele model. Importantly, the fact that Eq. (8.43) is rooted in Lindblad theory ensures that it respects total-probability conservation. Energy and spin are only conserved locally and dissipated at the boundary.

8.3 Macroscopic Edge Dynamics in Ribbon Segments with Lindblad Boundaries

In the following sections, we aim to simulate the edge dynamics of a macroscopic Kane–Mele model, i.e. a half-space model with a single physical zigzag edge. In order to achieve this, we consider a finite Kane–Mele tile with selective Lindblad bath coupling at its boundary, as described above. Specifically, we base our calculations on a *ribbon-segment geometry*. This is illustrated in Fig. 8.3, where the lower zigzag edge represents the physical edge and the remaining three edges are coupled to the Lindblad bath at the grey coloured sites. We call this tile geometry a ribbon segment geometry because its extent in the physical (zigzag) x -direction is much greater than that in the perpendicular (armchair) y -direction. The choice of a ribbon segment over a more symmetric tile is motivated by numerical efficiency: since the numerical cost of simulating the full real-time dynamics scales as L^2 for large L , we have to limit our considerations to systems of about $500 \lesssim L \lesssim 600$ sites. A zigzag ribbon segment can then be used to maximise the length of the physical zigzag edge while keeping the overall system size L fixed. Concretely, a zigzag ribbon segment is a tile with a large aspect ratio $\mathcal{R} = \ell_x/\ell_y$, where ℓ_x and ℓ_y are

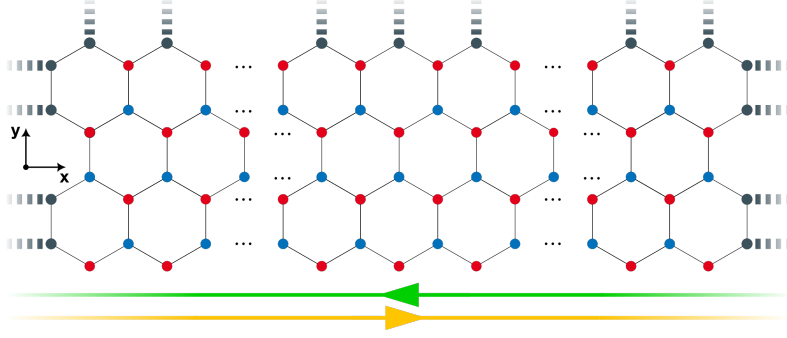


Figure 8.3: Sketch of a ribbon segment with Lindblad boundaries. Red (blue) dots mark sublattice A (B) sites. The grey dots indicate edge sites that are coupled to a Lindblad bath, shown as thick dashed grey lines; the uncoupled bottom zigzag edge corresponds to the physical edge of a half-space model. The green (yellow) lines at the bottom indicate the spin-up (spin-down) polarised helical edge modes, while the green (yellow) arrow heads signify the direction of spin-up (spin-down) transport. Reproduced with minor modifications from Ref. [RQ2].

the numbers of unit cells in the x - and y -directions. Since there are two sites per unit cell, the total number of lattice sites is given by $L = 2\ell_x\ell_y$. Accordingly, we obtain minimal and maximal aspect ratios of $\mathcal{R}_{\min} = 2/L$ and $\mathcal{R}_{\max} = L/2$ for fixed L . While the maximal aspect ratio $\mathcal{R}_{\max}$ does yield the longest possible physical zigzag edge, it cannot be used to simulate the edge dynamics of a half-space model, as the resulting system effectively reduces to a one-dimensional isolated zigzag edge along the x -direction with no perpendicular bulk dimension. In order to simulate a half-space model we therefore need $\ell_y > 1$. In fact, the minimal ℓ_y required for half-space simulations is determined by the interaction between the helical edge modes on opposite zigzag edges: if the ribbon segment is too narrow, these modes overlap, hybridise and gap out. Moreover, a significant overlap between opposite zigzag edge states would allow the Lindblad bath – coupled to the top zigzag edge – to influence dynamics at the physical (bottom) zigzag edge, cf. Fig. 8.3. We have determined that ribbon widths of about four unit cells along the armchair y -direction are typically sufficient to suppress hybridisation between the zigzag edge states in the topological phase. A total site count of $500 \lesssim L \lesssim 600$ therefore allows ribbon segment lengths of approximately $62.5 \lesssim \ell_x \lesssim 75$ unit cells along the zigzag x -direction. In practice, we choose ℓ_x to be half-integer. The extra half unit cell in the x -direction ensures that the zigzag edges always terminate on B -sublattice sites, cf. Fig. 8.3. Since the zigzag edge states are predominantly localised on the A -sublattice sites, this configuration suppresses direct interactions between the Lindblad bath and the helical edge modes at the two corners (bottom left and right corners in Fig. 8.3) where the physical zigzag edge meets the adjacent Lindblad armchair edges.

8.4 Spin Injection and Helical Transport

We begin by verifying the existence of helical spin transport along the zigzag edge of a topological Kane–Mele model. To this end, we consider a hybrid system

$$H = H_{\text{KM}} + H_{\text{I}} \quad (8.45)$$

of an electronic Kane–Mele model H_{KM} as given in Eq. (8.1) and a local magnetic injection term H_{I} as specified in Eq. (8.36). For the Kane–Mele model, we consider a ribbon segment of $L = 572$ sites with $\ell_x = 71.5$ unit cells along the zigzag direction and $\ell_y = 4$ unit cells in the armchair direction. Figure 8.3 provides an illustration of this setup. As discussed above, the lower zigzag edge is the physical edge, while the remaining three edges are coupled to a Lindblad bath to dissipate excitations. While increasing the system size generally reduces finite-size artefacts, we have confirmed that the real-time dynamics remain robust under both uniform scaling and more substantial modifications of the system geometry, such as changes in aspect ratio $\mathcal{R}$. This robustness implies that we are effectively considering the edge dynamics

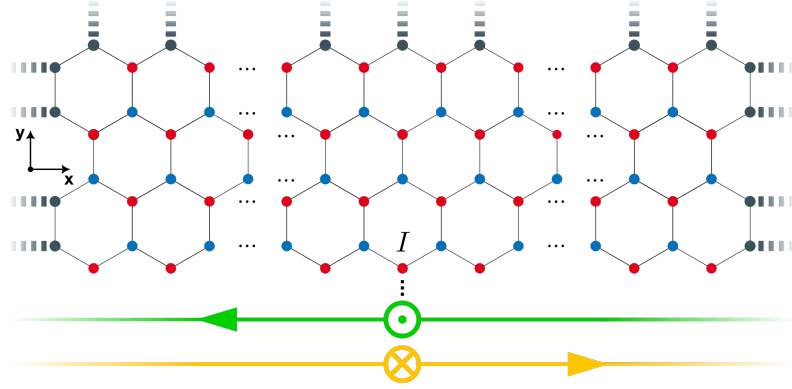


Figure 8.4: Sketch of a ribbon segment. Red (blue) dots mark the sites of sublattice A (B). The grey dots indicate edge sites that are coupled to a Lindblad bath, shown as thick dashed grey lines; the uncoupled bottom zigzag edge corresponds to the physical edge of a half-space model. The green $\odot$ (yellow $\otimes$) symbol represents an out-of-plane magnetic field $\mathbf{B}_I$ in positive (negative) z -direction coupled (dotted black line) to an “injection site” I of the physical zigzag edge. Here, I is chosen as the central site of the physical zigzag edge. A finite $\mathbf{B}_I = +|B_I|\mathbf{e}_z$ ($\mathbf{B}_I = -|B_I|\mathbf{e}_z$) induces a finite spin-up (spin-down) z -polarised spin density. The excitation spin-up (spin-down) excitation propagates to the left (right) once the injection field is switched off, cf. green (yellow) arrows at the bottom. Reproduced with minor modifications from Ref. [RQ2].

of a Kane–Mele half-space model. For the magnetic injection term H_I we use a z -polarised local magnetic field $\mathbf{B}_I = \pm B_I \mathbf{e}_z$. The injection site I is chosen to be the central site of the physical zigzag edge. The complete setup is sketched in Fig. 8.4, with the injection site labelled by I . The green $\odot$ (yellow $\otimes$) symbols indicate an injection field pointing along the positive (negative) z -direction, respectively. The injection fields in positive (negative) z -direction are shown in green (yellow) colour because they inject spin-up (spin-down) density⁶ into the system, which is expected to be transported by the corresponding spin-up (spin-down) polarised edge modes of the pristine Kane–Mele model represented by the green (yellow) lines in Fig. 8.3. This expectation is based on the following observations. Recall that the intrinsic SOC of the pure Kane–Mele model H_{KM} reduces its $\text{SU}(2)$ spin symmetry to a $\text{U}(1)$ symmetry around the z -axis. Despite this, the Kane–Mele ground state from Eq. (8.13) is an unpolarised Fermi sea of the form, $|\text{GS}\rangle = \prod_{\mathbf{k}} d_{\mathbf{k}\uparrow}^\dagger d_{\mathbf{k}\downarrow}^\dagger |0\rangle$ and hence a non-degenerate $\text{SU}(2)$ -symmetric spin singlet. Any additional local magnetic field $\mathbf{B}_I$ then reduces the ground state $\text{SU}(2)$ symmetry to a $\text{U}(1)$ symmetry, although the former invariance under $\text{SU}(2)$ ensures that the ground-state energy is independent of the direction of the field. With the choice $\mathbf{B}_I = \pm B_I \mathbf{e}_z$ that we use here, we therefore have a $\text{U}(1)$ spin-rotation invariance around the z -axis for both the Hamiltonian and its ground state. A field in the positive (negative) z -direction then induces a spin-up (spin-down) polarisation of the local magnetic moments in the vicinity of the injection site I . Since the helical boundary modes dominate the local density of states on the boundary sites, we expect this polarisation cloud to be predominantly supported by and propagate within the spin-momentum-locked edge states.

To test this numerically, we fix the Kane–Mele model parameters as $t_{\text{hop}} = 1$, $t_{\text{SO}} = 0.05$, $\Delta E = 0.3$, and choose V as defined in Eq. (8.12) for the topologically trivial and non-trivial cases, respectively. The magnetic field strength is set to $B_I = 1$, and the Lindblad parameter to $\gamma_0 = 0.2$. In the topologically trivial configuration, the resulting ground state has a local magnetic moment of $\langle s_I^z(t=0) \rangle \approx 0.16$ at the injection site I . This ground state polarisation is shown as the $t = 0$ polarisation in panel (a) of Fig. 8.5, where the upper-half (lower-half) purple data indicates the ground state polarisation induced by a magnetic field pointing along the positive (negative) z -direction. Panel (b) of Fig. 8.5 shows the same data for the topologically non-trivial configuration, where the ground state develops a significantly larger local magnetic moment of $\langle s_I^z(t=0) \rangle \approx 0.33$ at the injection site I . Note that the local moment is not fully polarised ($\langle s_I^z(t=0) \rangle < 0.5$) in either case. The key difference between the trivial and non-trivial setup shown in Fig. 8.5 is the dynamics triggered by the magnetic field injection at $t = 0$. These dynamics are visualised through snapshots of the local spin-polarisations $\langle s_i^z(t) \rangle$ at selected instants of times.

⁶This is due to the antiferromagnetic coupling of H_I from Eq. (8.36).

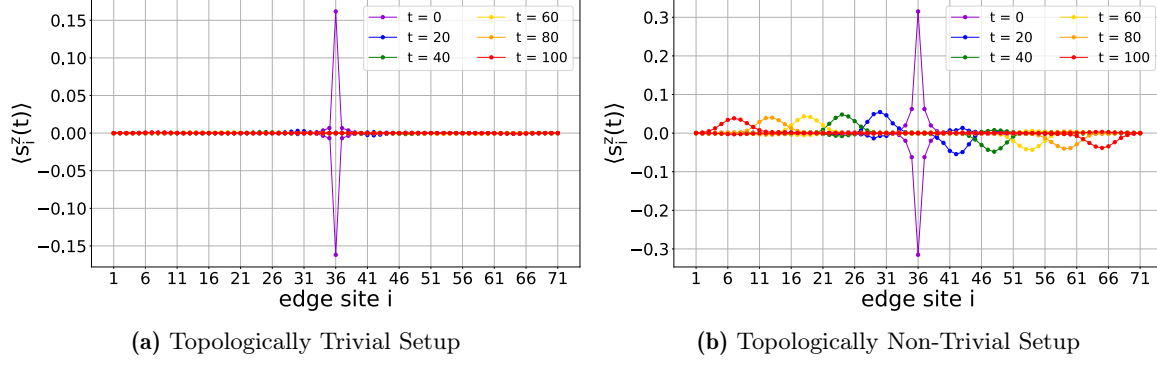


Figure 8.5: Spatial distribution of spin-polarisations $\langle s_i^z(t) \rangle$ at selected instants of time (colours). At $t = 0$, the spin-injection field B_I at edge site $I = 36$ is abruptly switched off, releasing the excitation; A -sublattice edge sites are numbered from $i = 1$ to 71 . The upper halves show a spin-up injection with $\langle s_i^z(t) \rangle > 0$, while the lower halves show a spin-down injection with $\langle s_i^z(t) \rangle < 0$. The time unit is set by $t_{\text{hop}} = 1$ and $\hbar \equiv 1$. The remaining parameters are $t_{\text{SO}} = 0.05$, $\Delta E = 0.3$ (as in Fig. 8.1), V as obtained from Eq. (8.12) with positive (a) and negative (b) sign, respectively. The initial field strength is $B_I = 1$, and the Lindblad parameter $\gamma_0 = 0.2$. Adapted with minor modifications from Ref. [RQ2].

In the topologically trivial case, shown in Fig. 8.5 (a), the polarisation cloud is tightly confined to the injection site I at $t = 0$. Once the magnetic injection field is removed, it shows little to no propagation along the edge. Instead, the polarisation cloud diminishes rapidly and essentially vanishes after about 10 inverse hoppings. This quick decay is due to the fact that almost the entire weight is immediately dissipated into the bulk of the system. The dynamics are almost perfectly symmetric with respect to the spin orientation: spin-up (upper half of the figure) and spin-down (lower half of the figure) injections lead to identical outcomes after only $t = 40$ inverse hoppings.

In contrast, the topologically non-trivial setup, shown in Fig. 8.5 (b), exhibits pronounced and highly asymmetric edge dynamics. At $t = 0$, the polarisation cloud is already spread over about five sites around the injection site I . Following the removal of the magnetic injection field, the excitation does not decay in place but propagates along the edge, broadening over time until it extends across roughly ten sites after $t = 100$ inverse hoppings. Most importantly, the dynamics are no longer symmetric with respect to the spin orientation. While the time-dependent shapes of the spin-up and spin-down polarisation profiles are the same, the spin-up excitation predominantly propagates to the left, whereas the spin-down counterpart propagates mostly to the right along the zigzag chain. This is clear evidence of the spin-momentum locking in the topologically non-trivial state. It is also worth noting that, despite the persistent and directional edge propagation, about half of the initial spin weight is lost within the first 10 inverse hoppings. This is attributed to the fact that about a half of the injected spin density is carried by bulk states and thus quickly dissipated into the bulk.

The data for the topologically non-trivial state in panel (b) of Fig. 8.5 also demonstrate that, after the initial dissipation of spin density into bulk states, the total weight of the remaining polarisation cloud remains nearly constant over time. This is consistent with our expectations, since the residual spin density is almost exclusively carried by the corresponding helical edge modes, which support lossless spin transport. Furthermore, the propagation velocity of the spin density “wave packets” is approximately $v \approx 0.29$ sites per inverse hopping, as determined from the peak positions. This is in excellent agreement with the Fermi velocity $v_F \approx 0.285$ of the helical edge states, as mentioned in the discussion of Fig. 8.1.

Finally, we note that there is a low-weight spin-down (spin-up) density peak around site $i = 29$ ($i = 43$) in the spin-up (spin-down) case at $t = 20$. This peak is located to the left (right) of the injection site $I = 36$ and continues propagating in that direction at later times, as suggested by the wave packets at $t = 40$. This reversal of the expected helical propagation direction indicates a partial breakdown of spin-momentum locking. Importantly, this is not due to an explicit violation of TRS *during* the evolution, but rather a remnant of the locally TRS-breaking *initial state preparation*. In fact, the TRS-transformed dynamics of the spin-up dynamics – involving a simultaneous reversal of time (reflection across the y -axis of the figure) and spin (reflection across the x -axis of the figure) – precisely match the observed spin-down dynamics, confirming that the dynamics of the system are governed by a TRS invariant Hamiltonian.

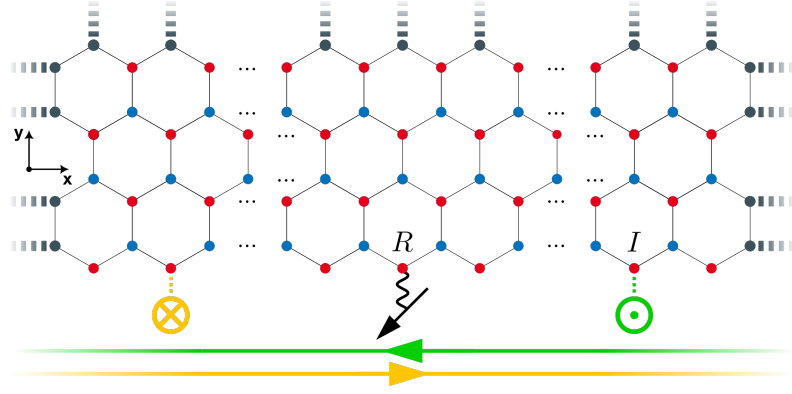


Figure 8.6: Same as Fig. 8.4, but with an additional classical read-out spin $\mathbf{S}_R$, shown as a black arrow, which is antiferromagnetically coupled to the local spin at the central site R of the physical zigzag edge. The coupling is illustrated as the wavy black line. Yellow $\otimes$ (green $\odot$) symbols mark “proper” (cf. main text) locations I of spin-up (spin-down) injection fields $\mathbf{B}_I = +B_I \mathbf{e}_z$ ($\mathbf{B}_I = -B_I \mathbf{e}_z$) half way between R and the bottom left (bottom right) of the corners. The proper injection locations ensure that the spin-momentum locked propagation of the injected spin-density in the helical boundary modes (green and yellow lines at the bottom) predominantly moves it towards the read-out spin at R . Reproduced with minor modifications from Ref. [RQ2].

8.5 Helical Transport in the Presence of the Read-Out Spin

Having demonstrated helical spin-density transport in the pristine Kane–Mele half-space model, we now include the classical read-out spin and turn to the full dynamical setup described by Eq. (8.37). This is sketched in in Fig. 8.6. In addition to the spin-up (spin-down) injection field – once more indicated by the green $\odot$ (yellow $\otimes$) symbols – a classical read-out spin $\mathbf{S}_R$ (black arrow) with $S_R \equiv |\mathbf{S}_R| = 1/2$ is exchange coupled to the central site R of the physical zigzag edge. We choose an antiferromagnetic exchange coupling strength of $J = 2$ for the classical read-out spin, such that $J|\mathbf{S}_R| = 1$. The spin-up (spin-down) injection field is again aligned along the positive (negative) z -axis, i.e. $\mathbf{B}_I = \pm B_I \mathbf{e}_z$, and positioned at a site I half way between R and the bottom right (left) corner of the ribbon segment. This placement of the spin-up (spin-down) injection field ensures that the helical propagation of the injected spin density will direct it towards the read-out spin, rather than away from it. Furthermore, we initialise the classical read-out spin in-plane, e.g. $\mathbf{S}_R = S_R \mathbf{e}_x$, as this maximises the torque exerted by the propagating spin excitation generated by $\mathbf{B}_I$. The initial state of the electron system is prepared as the ground state for fixed orientations of $\mathbf{B}_I$ and $\mathbf{S}_R$.

Before we move on to the discussion of the read-out spin dynamics, we examine the dynamics of the injected spin-polarisation cloud in the presence of the read-out impurity spin. To this end, we consider a Kane–Mele ribbon segment of the same size ($L = 572$) and dimensions ($\ell_x = 71.5$ and $\ell_y = 4$) as in the previous section, and couple an x -polarised classical read-out spin $\mathbf{S}_R = S_R \mathbf{e}_x$ to the former position of the injection field $\mathbf{B}_I$ at the central site $R = 36$ of the physical zigzag edge. Depending on whether we wish to study a spin-down or a spin-up polarisation cloud, we then add a magnetic injection field $\mathbf{B}_I = -B_I \mathbf{e}_z$ or $\mathbf{B}_I = +B_I \mathbf{e}_z$ to the injection site $I = 18$ or $I = 54$ half way between the central read-out site R and the bottom left or bottom right corner of the ribbon segment. Once the electronic system is initialised for a chosen configuration, we switch off the injection field to release the accumulated spin density, allowing it to propagate freely and interact with the read-out spin.

Figure 8.7 visualises this process for a spin-down density injection. At $t = 0$, the polarisation cloud is localised near site $I = 18$, located to the left of the read-out spin. Its profile is nearly identical to that of the spin-down polarisation cloud formed in the absence of a read-out spin, cf. Fig. 8.5, indicating that the spin injection process itself is largely unaffected by the reduced distance to the nearest armchair edge and the presence of the read-out spin. For $t > 0$, the spin density propagates mainly to the right, which is consistent with the spin-momentum locking we have seen earlier in the pristine Kane–Mele model. This behaviour persists despite the fact that the presence of the read-out spin formally breaks the TRS

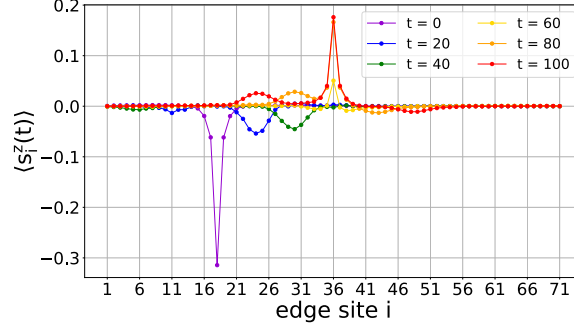


Figure 8.7: Similar to the spin-down injection in Fig. 8.5, this setup includes the classical read-out spin, which is antiferromagnetically coupled to edge site $R = 36$ with exchange coupling strength $J = 2$. Accordingly, the injection site is relocated to $I = 18$, to the left of R , such that helical propagation directs the spin-down injection towards the read-out spin. At $t = 0$, the electron system is in its ground state for fixed $\mathbf{B}_I(t = 0) = -B_I \mathbf{e}_z$ and $\mathbf{S}_R = S_R \mathbf{e}_x$. For $t > 0$, the injection field is switched off, i.e. $\mathbf{B}_I(t > 0) = 0$. Other parameters: see Fig. 8.5. Adapted with minor modifications from Ref. [RQ2].

invariance of the total Hamiltonian, reaffirming that local TRS breaking is indeed a valid notion [150]. Note that the spin density profiles at $t = 20$ and $t = 40$ suggest that the polarisation cloud experiences no significant perturbations due to the read-out spin prior to their direct interaction after about $t \approx 60$ inverse hoppings. Following this interaction, the spin-down excitation scatters strongly off the read-out spin. As a result, part of it is reflected as a backward-propagating spin-up excitation, another part becomes bound by the deflected read-out spin, and a small fraction is transmitted through the impurity, continuing in the spin-down channel. Combined, these observations corroborate the view that the TRS protected helical boundary modes remain intact away from the vicinity of the locally TRS-breaking read-out spin. Moreover, the spin-up polarisation cloud bound by the read-out spin shows minimal change for $t \geq 80$ (compare orange and red data points close to $R = 36$). This suggests that the read-out spin dynamics responsible for “trapping” the spin-up density have effectively come to an end by $t = 100$. The fact that these dynamics stabilise with a finite spin-up polarisation bound to the impurity site implies that the read-out spin has acquired a finite z -component via the torque exerted by the polarisation cloud during the scattering process.

8.6 Helical Control over a Read-Out Spin

This interpretation is in fact reinforced by the classical read-out spin dynamics. To see this, we contrast the read-out spin dynamics following a “proper” injection process, where the injection field is positioned (as described earlier and shown in Fig. 8.7) such that the helical propagation moves the injected spin *towards* the read-out spin, with those following an “improper” injection process, where the injection field is instead positioned such that helical propagation moves the spin injection *away from* the read-out spin. Figure 8.8 compares the read-out spin dynamics resulting from “proper” and “improper” injection processes by overlaying the trajectories traced out by the read-out spin tip throughout the dynamics. These are visualised as curves A-D on the two-sphere $S^2_{1/2}$ with radius $S_R = 1/2$. Specifically, curves A and B describe the read-out spin dynamics induced by “proper” and “improper” spin-down injections, while curves C and D illustrate the same for “proper” and “improper” spin-up injections. As expected, the “proper” spin-up (spin-down) injections (cf. high-weight polarisation clouds in Fig. 8.7) cause a substantial deflection of the read-out spin towards the positive (negative) z -direction. The fact that the helically suppressed “improper” spin injections (cf. low-weight polarisation clouds in Fig. 8.7) provoke a much weaker response supports the interpretation that the read-out spin deflection is *primarily driven* by the torque of the propagating spin polarisation clouds.

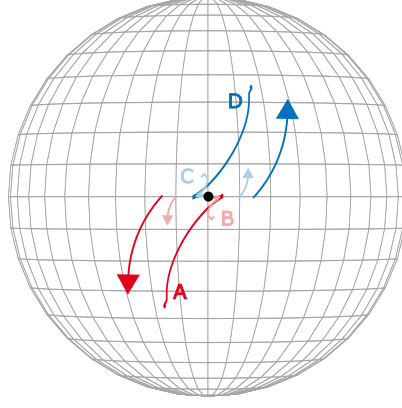


Figure 8.8: Trajectories of the read-out spin tip on the $\mathbb{S}^2$ configuration space is shown for processes A-D, each initiated by a static spin injection. The injection field $\mathbf{B}_I$ is switched off at $t = 0$, and the system’s state is propagated until $t = 150$. The read-out spin $\mathbf{S}_R$ is located at the center $R = 36$ of the physical edge and initially aligned along the $+x$ -direction (black dot). A (deep red): “proper” spin-down injection at site $I = 18$ to the left of site R . B (faint red): “improper” spin-down injection at site $I = 54$ to the right of site R . C and D (deep and faint blue): “proper” and “improper” spin-up injection at $I = 54$ and $I = 18$, respectively. Parameters as in Fig. 8.7. The arrows indicate the direction of the spin-tip deflection for each process. Adapted with minor modifications from Ref. [RQ2].

However, not all of the observed read-out spin dynamics can be attributed to the propagating spin densities. In particular, all of the trajectories A-D exhibit a non-trivial *in-plane dynamics* that sets in immediately after the injection field is switched off. The reason for this is that the initial state of the electron system, i.e. ground-state Fermi sea for fixed orientations of $\mathbf{B}_I$ and $\mathbf{S}_R$, is “stressed”, meaning it differs from the ground state of the total system. The latter also minimises the total energy with respect to the orientation of $\mathbf{S}_R$. Specifically, the total system ground state is realised when both, $\mathbf{S}_R$ and $\mathbf{B}_I$, lie in the xy -plane, enclosing a relative azimuthal angle $\Delta\phi(r)$ that depends on the inter-impurity distance r . According to Ref. [193], this angle is given by $\Delta\phi(r) = \pi - \alpha(r)$, where $\alpha(r) = 2E_F r / v_F$ with the inter-impurity distance r , the Fermi energy E_F , and the Fermi velocity v_F . For $E_F = 0$, the relative angle $\Delta\phi(r)$ becomes independent of the inter-impurity distance, and the impurities align antiferromagnetically with $\Delta\phi = \pi$ for all impurity separations r . However, $E_F = 0$ is only realised exactly in the thermodynamic limit $L \rightarrow \infty$. In finite systems, like the one considered here, E_F typically deviates slightly from zero, resulting in a small but finite distance-dependent deviation $\alpha(r) \neq 0$ from the antiferromagnetic relative angle $\Delta\phi = \pi$. Given that the injection field is aligned along the z -direction in our setup, the initial state always constitutes a slightly excited state, whose local magnetic moments do not align⁷ with those of the electronic ground state at $\mathbf{B}_I = 0$. The resulting magnetic stress initiates spin relaxation dynamics, which begin immediately after the injection field is removed, exerting a finite torque on $\mathbf{S}_R$. However, this relaxation dynamics is rather weak. As shown in Fig. 8.8, the spin drifts only a few degrees on the two-sphere, staying mostly within the xy -plane. A qualitative change of the dynamics only occurs when the injected polarisation cloud arrives after about $t \approx 50$ inverse hoppings. As previously mentioned, the extent of this excitation-induced dynamics depends on the weight of respective polarisation cloud: it is large for the helically propagating “proper” injections and small for the helically suppressed “improper” injections. In both cases, the read-out spin dynamics slow to a halt after a total of about $t \approx 150$ inverse hoppings. Finally, it is worth mentioning that the initial electron ground states of setups with spin-up and spin-down injection fields are stressed with *opposite helicity*. This results in a relative sign change of the torque on $\mathbf{S}_R$, which explains the perfect symmetry between the curves A (B) and D (C) in Fig. 8.8.

⁷The extent of this discrepancy depends on the distance between I and R .

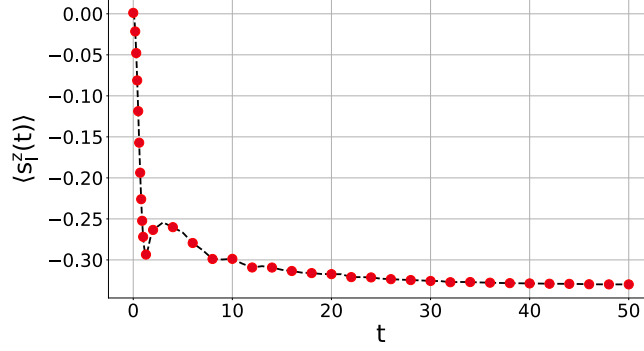


Figure 8.9: Real-time evolution of the z -component $\langle s_i^z(t) \rangle$ of the local magnetic moment at site $I = 18$ on the physical zigzag edge during a spin-down injection driven by an injection field $\mathbf{B}_I$ applied in the negative z direction for $t_{\text{inj}} = 50$ inverse hoppings. Parameters correspond to the topologically non-trivial case shown in Fig. 8.7. Adapted with minor modifications from Ref. [RQ2].

8.7 Iterated Injection Protocol: Building a Helical Spin Switch

The fact that the read-out spin dynamics comes to a halt once the injected polarisation cloud has passed by shows that the system is, at least locally, in a state close to its ground state. We can use this near-equilibrium state to initialise a subsequent dynamical process. This possibility raises a number of natural questions: for instance, can we undo the deflection of the read-out spin by a suitable subsequent process? Moreover, can we achieve a complete switching $\pm S_R \mathbf{e}_z \rightarrow \mp S_R \mathbf{e}_z$ between the north and south pole orientations of the read-out spin by iterating the sequence studied so far? If so, is it possible to undo that process as well?

To start the discussion, we consider a single additional “basic injection-pump” (BIP) process. This BIP process involves (i) a spin injection, which, in contrast to the injection processes discussed so far, must be treated fully dynamically, and (ii) the subsequent pumping of the read-out spin dynamics driven by the propagating spin injection. To distinguish the dynamically implemented spin injections from the previously discussed spin injections via initial state preparation, we will henceforth refer to the former as dynamic injection processes and the latter as static injection processes. A dynamic spin injection starts from an arbitrary near-equilibrium state of the total system, i.e. a state where the electron system is almost in the ground state for a given *arbitrary* orientation of the read-out spin $\mathbf{S}_R$. One could, in principle, conceive a stricter construction scheme where the total system is in perfect equilibrium, i.e. where the electron system is in the exact ground state for a given $\mathbf{S}_R$, rather than just close to it. However, this is not practical for our purpose: since we are aiming to develop a workable dynamical protocol of iterated BIP processes, we cannot afford to wait for full thermalisation after every dynamic injection. An isolated dynamic spin injection then proceeds as follows. At $t = 0$, the spin-injection field $\mathbf{B}_I = \pm B_I \mathbf{e}_z$ is switched on, causing an accumulation of local magnetic moment around the injection site I . The built-up magnetic moment forms a polarisation cloud, which continues to develop until the injection field is switched off after a specified “injection” time t_{inj} . Unsurprisingly, the magnitude of the final polarisation depends on t_{inj} . This is demonstrated in Fig. 8.9, which shows the temporal evolution of the local magnetic moment $\langle s_i^z(t) \rangle$ at the injection site I during a dynamic spin-down injection ($\mathbf{B}_I = -B_I \mathbf{e}_z$) performed on the electronic ground state for a fixed classical read-out spin oriented along the positive x -direction. We observe that the formation of a local magnetic moment parallel to the injection field at I happens very quickly. In fact, it takes only a few ($t \approx 2$) inverse hoppings for $|\langle s_i^z(t) \rangle|$ to reach a value ($|\langle s_i^z(t) \rangle| \approx 0.3$) close to its eventual saturation limit ($|\langle s_i^z(t) \rangle| \approx 0.33$). After a brief rebound ($2 \lesssim t \lesssim 5$), the local magnetic moment continues to converge monotonically ($t > 5$), approaching its limit within a total of 40 to 50 inverse hoppings. The injection time t_{inj} can then be determined as a point in time at which $|\langle s_i^z(t) \rangle|$ has converged to some sufficient degree. Here, the local magnetic moment essentially converges within about $t_{\text{inj}} \equiv 50$ inverse hoppings, after which we get $\langle s_i^z(t_{\text{pump}}) \rangle \approx -0.33$, showing that the field strength of $B_I = 1$ is not sufficient to fully polarise the magnetic moment. We will stick to $t_{\text{inj}} = 50$ in the following. Note that the results are practically independent of the choice of I , provided that the distance to the corners and to R is large enough. A few sites have proven to be sufficient. Once

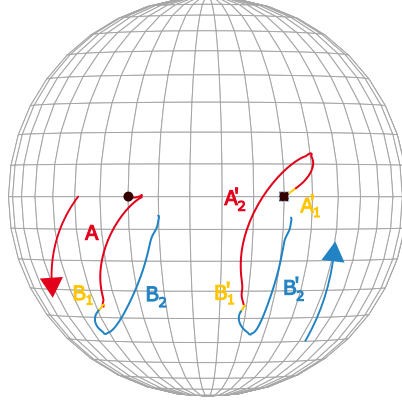


Figure 8.10: *Left:* Trajectory of the read-out spin tip on the S^2 configuration space. Process A (red), as in Fig. 8.8, is a static spin-down injection at site $I = 18$ with read-out at $R = 36$. The injection field is switched off at $t = 0$, and the system evolves until $t_{\text{pump}} = 150$. Process B is a spin-up BIP sequence with injection at $I = 54$ and read-out at $R = 36$, consisting of B_1 (yellow), a dynamic injection over $t_{\text{inj}} = 50$ starting from the final state of A, and B_2 (blue), the subsequent release and propagation of the spin-up density, including scattering at $\mathcal{S}_R$, up to $t_{\text{pump}} = 150$. The black dot marks the initial orientation of $\mathcal{S}_R(t = 0) = S_R \mathbf{e}_x$. *Right:* the same process but starting, for better visibility, from $\mathcal{S}_R(t = 0) = S_R(\mathbf{e}_x + \mathbf{e}_y)/\sqrt{2}$ (black square), and with A replaced by A' consisting of a proper dynamic spin-down injection A'_1 and a pump part A'_2 . B'_1 and B'_2 are the same as B_1 and B_2 but starting from the final state of process A' . Adapted with minor modifications from Ref. [RQ2].

the injection field $\mathbf{B}_I$ is switched off at $t = t_{\text{inj}}$, the formation of the spin polarisation cloud stops and it is released into the electronic system. If the position of the injection site I relative to that of the read-out site R is chosen properly, the polarisation cloud propagates towards and scatters from the read-out spin, pumping its dynamics as described before. The backscattered and transmitted parts of the polarisation cloud are eventually dissipated into the Lindblad bath, simulating total-energy-conserving dissipation in a macroscopically large sample. As a result, the pump dynamics of the read-out spin come to an end after a finite time t_{pump} . In our case, a pump duration of $t_{\text{pump}} = 150$ inverse hoppings was found to be sufficient, see the discussion in Secs. 8.5 and 8.6. Generally, a sensible lower bound for t_{pump} is given by the distance between sites I and R divided by the Fermi velocity v_F of the helical edge states. The BIP process can be terminated after $t \approx t_{\text{inj}} + t_{\text{pump}}$ inverse hoppings and another BIP process may follow.

In order to determine whether a BIP process can undo the proper read-out spin deflection shown in Fig. 8.8, we consider a concatenated process A+B consisting of a proper static spin-down injection A (cf. Sec. 8.6) and a proper spin-up BIP process $B = B_1 + B_2$ involving a (proper) dynamic spin-up injection B_1 and the subsequent pump dynamics B_2 . The combined process A+B is illustrated on the left side of Fig. 8.10. It starts at time $t = 0$ from an initial state given by the electronic ground state for a fixed x -polarised (black dot) read-out spin at site $R = 36$ and a fixed spin-down ($\mathbf{B}_I = -B_I \mathbf{e}_z$) injection field at site $I = 18$ (static spin injection setup). As described before, the deactivation of the injection field triggers system dynamics, during which the read-out spin moves towards the south pole of the two-sphere (red line) – this is the same as process A from Fig. 8.8. After about $t = 150$ inverse hoppings, the dynamics of $\mathcal{S}_R$ have essentially ceased and the physical (zigzag) edge of the electronic system has relaxed to a “final” state close to the (local) ground state. Immediately after the static spin-down injection A is terminated at $t = 150$, the spin-up BIP process begins with a dynamic spin-up injection B_1 (yellow line) applied to this nearly relaxed final state of the previous process. The dynamic spin-up injection B_1 lasts for $t_{\text{inj}} = 50$ inverse hoppings (cf. Fig. 8.9), and Fig. 8.10 shows that it has only a minimal effect on the orientation of the read-out spin. Once the dynamic injection is completed, the accumulated spin-up polarisation cloud is released, and its helical propagation drives the pump dynamics B_2 (blue line) of the read-out spin. These last for another $t_{\text{pump}} = 150$ inverse hoppings. After a brief continuation of the downward motion in the early part of B_2 , the spin-up polarisation cloud scatters and drives the read-out spin back towards the north pole of the two-sphere. Although it does not quite return to its original position $\mathcal{S}_R = S_R \mathbf{e}_x$ at the equator, we still observe that the spin-up BIP process B largely reverses the preceding static spin-down process A. There are two main reasons why the reversal is incomplete.

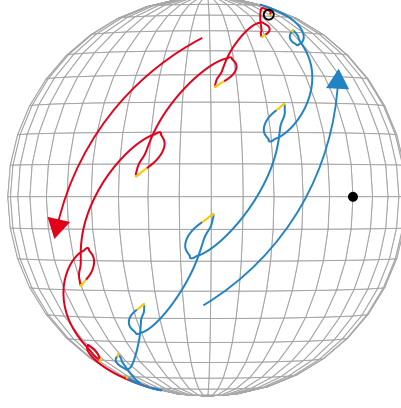


Figure 8.11: Trajectory of the read-out spin tip on the S^2 configuration space throughout a process consisting of 6 spin-down followed by 7 spin-up BIP processes starting from a state (open black dot) close to the north pole, $\eta = 0.05$, as discussed in the text. The filled black dot marks the position $S_R \mathbf{e}_x$ for reference. The yellow lines represent spin injections (50 inverse hoppings). The red (blue) lines indicate proper spin-down (spin-up) pump (150 inverse hoppings). Parameters as in Figs. 8.5 and 8.7. Adapted with minor modifications from Ref. [RQ2].

The first one is that the processes A and B are not *constructed* as exact inverses of one another: while A starts from a static injection, B is based on a dynamic one. To quantify the impact of this discrepancy, we repeated the whole cycle using two dynamic spin injection protocols $A' = A'_1 + A'_2$ (spin-down BIP) and $B' = B'_1 + B'_2$ (spin-up BIP), both of which are executed as described above ($t_{\text{inj}} = 50$, $t_{\text{pump}} = 150$). This is shown on the right side of Fig. 8.10. For a clearer visual representation, we chose a slightly rotated initial orientation $\mathbf{S}_R = S_R(\mathbf{e}_x + \mathbf{e}_y)/\sqrt{2}$ (black square) for the read-out spin. The concatenated process $A' + B'$ causes a significantly smaller deviation from its initial state than $A + B$, allowing us isolate which part of the deviation is caused by the combination of different injection mechanisms.

The second reason for the incomplete inversion is that in all of the above protocols, the sub-processes are terminated before the system has fully relaxed. This affects both the asymmetric ($A + B$) and the symmetric ($A' + B'$) inversion protocol alike, and explains why the latter still fails to close, even though the two sub-processes are exact inverses of one another. If each sub-process X was extended until the electronic system is fully relaxed to its ground state, the final state of X would be *equivalent* to its initial state, and one could perfectly undo A' with B' . Here, equivalence of the initial and final states of sub-process X means that they differ only by a rotation $R(\mathbf{e}, \varphi)$ through some angle φ around an axis $\mathbf{n}_e$. Concretely, $\mathbf{n}_e$ is defined by the normal vector $\mathbf{e}$ to the plane spanned by the initial and final orientations of the classical read-out spin. Note that the electronic ground states $|\mathbf{S}_R\rangle$ and $|\mathbf{S}'_R\rangle$, corresponding to two different orientations $\mathbf{S}_R$ and $\mathbf{S}'_R$ of the classical read-out spin, have the same energy and are related by $|\mathbf{S}'_R\rangle = U(\mathbf{e}, \varphi)|\mathbf{S}_R\rangle$ where $U(\mathbf{e}, \varphi)$ is a unitary operator implementing $R(\mathbf{e}, \varphi)$. This is due to the fact that, as discussed in Sec. 8.4, the ground state of the pristine Kane–Mele model (without coupling to $\mathbf{S}_R$) is an SU(2)-invariant singlet, despite the SOC anisotropy.

A natural way to reduce the deviation caused by incomplete relaxation is to extend the pump duration of the BIP processes. We therefore conclude that an inversion of a proper spin-up (spin-down) BIP process by a proper spin-down (spin-up) BIP process is possible to arbitrary precision. This answers the first of the two questions we asked at the beginning of this section in the affirmative. Thus, we move on to the second question of whether it is possible to switch the read-out spin back and forth between the north and south pole orientations by suitable sequences of proper BIP processes. We begin by noting that incomplete relaxation of individual BIP processes is no longer a concern in this context, as it can always be compensated for by further BIP processes. That said, a new challenge arises as well. Namely, if the read-out spin is perfectly polarised along z -axis, i.e. if $\mathbf{S}_R = \pm S_R \mathbf{e}_z$, it remains unaffected by a passing spin-up or spin-down polarisation cloud, as the resulting torque $J\mathbf{S}_R \times \langle \mathbf{s}_R \rangle$ vanishes in this case. For this reason, we begin with an initial state, where the read-out spin is slightly tilted away from the z -axis. Specifically, we choose an initial orientation for $\mathbf{S}_R$ that is nudged towards the positive x -direction,

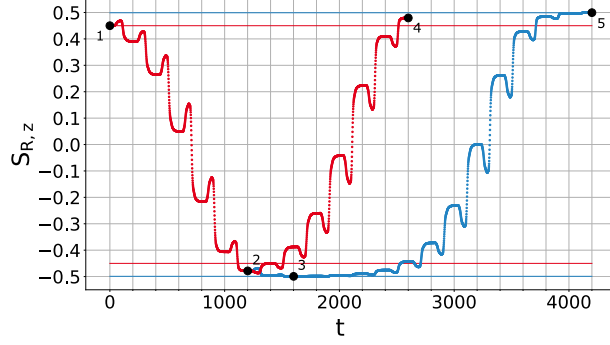


Figure 8.12: Time evolution of $S_{R,z}$ throughout full switching protocols. Red data points show data for the $\eta = 0.05$ switching protocol from Fig. 8.11. Labels 1, 2, 4 indicate start, reversal and termination points of the switching cycle. Horizontal red lines highlight threshold levels $\pm S_R(1 - \eta)$. Blue data points and lines use the same starting point (1) but $\eta = 0.001$ for the reversal and termination of the process, see labels 3 and 5. Vertical lines separate the different BIP processes, each lasting 200 inverse hoppings. Adapted with minor modifications from Ref. [RQ2].

writing

$$\mathbf{S}_R(t=0) = S_R(\sqrt{1 - (1 - \eta)^2} \mathbf{e}_x + (1 - \eta) \mathbf{e}_z), \quad (8.46)$$

where $\eta \in \mathbb{R}$ is a small parameter controlling the deviation of the read-out spin from the z -axis for $t = 0$. Figure 8.11 shows the trajectory traced by the tip of the classical read-out spin during an attempted full switching process. This process begins with six proper spin-down BIP processes (red and in-between yellow lines), which gradually reduce the z -component of $\mathbf{S}_R$ until $S_R^z(t) < -S_R(1 - \eta)$. At this point, $\mathbf{S}(t)$ is more closely aligned with $-S_R \mathbf{e}_z$ than $\mathbf{S}(0)$ was with $+S_R \mathbf{e}_z$. Thus, the parameter η also provides a natural termination criterion for the switching process. We find that within a 5% tolerance ($\eta = 0.05$), a full switching process can, in fact, be achieved by only 6 BIP processes. Each BIP process involves 50 inverse hoppings of proper spin injection and 150 inverse hoppings of free time evolution, during which the accumulated polarisation cloud propagates helically and pumps the dynamics of the read-out spin. Figure 8.11 also shows that, at the same tolerance level of $\eta = 0.05$, a comparable number of opposite BIP processes, in this case seven proper spin-up BIP processes (blue and in-between yellow lines), suffices to reverse the switching and return the read-out spin to an orientation close to its initial orientation $\mathbf{S}(0)$.

In order to quantify the impact of the tolerance level on the switching process, we make a comparison with a process using a stricter tolerance level. To this end, we overlay the time-dependent z -component $S_{R,z}(t)$ of $\mathbf{S}_R(t)$ for two switching processes of different tolerance levels. The red data in Fig. 8.12 corresponds to the $\eta = 0.05$ switching protocol discussed above. Each BIP process lasts for $t_{\text{inj}} + t_{\text{pump}} = 50 + 150 = 200$ inverse hoppings as indicated by the vertical lines. During the proper dynamic spin-up and spin-down injections, the read-out spin drifts mostly along a latitude, giving rise the plateaus extending over about $t_{\text{inj}} = 50$ inverse hoppings. The subsequent helical spin-down (spin-up) propagation then produces slow upwards (downwards) deflections which show up as moderate⁸ upwards (downwards) slopes extending over roughly $|R - I|/v_F \approx 62$ inverse hoppings after each plateau. The scattering of the spin-down (spin-up) polarisation cloud triggers a sharp downwards (upwards) deflection of the read-out spin. These pump dynamics show up as steep⁹ downwards (upwards) slopes. Together, the moderate upwards (downwards) slopes during the spin-down (spin-up) propagation and the steep downwards (upwards) slopes during the scattering form the characteristic bumps (dips) in the $S_{R,z}(t)$ trajectories. The blue curve in Fig. 8.12 shows a similar switching process but with $\eta = 0.001$. Concretely, the switching process starts at the $\eta = 0.05$ initial point and then proceeds according to the stricter tolerance level of $\eta = 0.001$. As a result, two additional (a total of eight) proper spin-down BIP processes are required to take the read-out spin from the $\eta = 0.05$ initial orientation to an $\eta = 0.001$ final orientation, where $S_{R,z}(t) < -S_R(1 - \eta)$ for $\eta = 0.001$. The reversion of the downward switching process takes still more effort, requiring a total of thirteen proper spin-up BIP processes to get from an $\eta = 0.001$ initial (down) orientation to an $\eta = 0.001$ final (up) orientation.

⁸The moderate deflection speed is evident from the tight clustering of time-equidistant data points along the slopes.

⁹The higher speed is reflected in the sparser spacing of time-equidistant data points along the slope.

9 – Sombrero Berry-Phases in BdG Vacuum Manifolds

An interacting microscopic model for superconductivity is described by a Hamiltonian H that is invariant under $U(1)$ gauge transformations corresponding to particle number conservation. The superconducting ground state of such a Hamiltonian spontaneously breaks this $U(1)$ symmetry, giving rise to a manifold of degenerate ground states. In the familiar image of the Sombrero-shaped energy landscape, these spontaneously broken ground states are parametrised by the phase of the superconducting order parameter, tracing the brim of the Sombrero.

Here, we do not address the spontaneously $U(1)$ -breaking ground states of interacting models. Instead, we focus on the family of ground states arising from the explicitly $U(1)$ -breaking BdG mean-field Hamiltonians of such models. Specifically, we consider the bundle of instantaneous BdG ground states (vacua) over $U(1) \simeq \mathbb{S}^1$, and formally determine the Berry phase acquired along closed paths in $\mathbb{S}^1$ using the geometric constructions outlined in Sec. 4. In reference to the manifold of spontaneously broken superconducting ground states, we refer to the resulting Berry phases as *Sombrero* Berry phases. This analysis is applied to the standard BCS model and the Kitaev chain model using the BdG formalism developed in Sec. 5.

9.1 Bogoliubov Diagonalisation of Reduced BdG Hamiltonians

Some BdG Hamiltonians may be reduced to the form

$$H = \sum_{\mathbf{k}} (c_{\mathbf{k}\alpha}^\dagger \ c_{-\mathbf{k}\beta}) \begin{pmatrix} \xi(\mathbf{k}) & \Delta(\mathbf{k}) \\ \Delta(\mathbf{k})^* & -\xi(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}\alpha} \\ c_{-\mathbf{k}\beta}^\dagger \end{pmatrix} =: \sum_{\mathbf{k}} \Psi(\mathbf{k})^\dagger h(\mathbf{k}) \Psi(\mathbf{k}), \quad (9.1)$$

where $\xi(\mathbf{k}) = \epsilon(\mathbf{k}) - \mu$ and $\Delta(\mathbf{k})$ are the single-particle energy dispersion relative to the chemical potential and the gap function. The indices $\alpha \neq \beta$ account for distinct internal degrees of freedom like (pseudo-)spin or orbitals. Equation (9.1) does not constitute a generic BdG Hamiltonian because the “Nambu” spinors $\Psi(\mathbf{k})$ contain no redundant single-particle information: the α -flavoured single particle states appear only as annihilation operators, while the β -flavoured ones appear only as creation operators. Still, Eq. (9.1) can be diagonalised by a Bogoliubov transformation,

$$U(\mathbf{k})^\dagger h(\mathbf{k}) U(\mathbf{k}) = \mathcal{E}(\mathbf{k}), \quad (9.2)$$

where

$$U(\mathbf{k}) = \begin{pmatrix} u(\mathbf{k}) & -v(\mathbf{k}) \\ v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} \quad (9.3)$$

is a unitary Bogoliubov matrix with elements

$$u(\mathbf{k}) = \sqrt{\frac{1}{2} \left(1 + \frac{\xi(\mathbf{k})}{E(\mathbf{k})} \right)} \quad \text{and} \quad v(\mathbf{k}) = \sqrt{\frac{1}{2} \left(1 - \frac{\xi(\mathbf{k})}{E(\mathbf{k})} \right)} \frac{\Delta(\mathbf{k})}{|\Delta(\mathbf{k})|} =: v_0(\mathbf{k}) e^{i\delta}, \quad (9.4)$$

and where

$$\mathcal{E}(\mathbf{k}) := \begin{pmatrix} E(\mathbf{k}) & 0 \\ 0 & -E(\mathbf{k}) \end{pmatrix} \quad (9.5)$$

is a real diagonal matrix of eigenenergies

$$E(\mathbf{k}) = \sqrt{\xi(\mathbf{k})^2 + |\Delta(\mathbf{k})|^2}. \quad (9.6)$$

The functions $u(\mathbf{k})$ and $v(\mathbf{k})$ from Eq. (9.4) satisfy

$$|u(\mathbf{k})|^2 + |v(\mathbf{k})|^2 = u(\mathbf{k})^2 + v_0(\mathbf{k})^2 = 1, \quad (9.7)$$

with $u(\mathbf{k}), v_0(\mathbf{k}) \in \mathbb{R}$ and $v(\mathbf{k}) \in \mathbb{C}$. Note that $v(\mathbf{k})$ carries the same complex phase as the superconducting gap $\Delta(\mathbf{k})$ and we wrote $v(\mathbf{k}) = v_0(\mathbf{k}) e^{i\delta}$ to tidy up the notation. A proof of Eq. (9.2) is given in App. A.11.

9.2 The Sombrero Berry-Phase of the BCS Ground State

The standard BCS chain Hamiltonian reads

$$H_{\text{BCS}} = \sum_{j,\alpha} \left(-t \left(c_{j\alpha}^\dagger c_{j+1\alpha} + c_{j+1\alpha}^\dagger c_{j\alpha} \right) - \mu c_{j\alpha}^\dagger c_{j\alpha} \right) - \sum_j \left(\Delta c_{j\downarrow}^\dagger c_{j\uparrow}^\dagger + \Delta^* c_{j\uparrow} c_{j\downarrow} \right). \quad (9.8)$$

Upon Fourier transforming the electronic field operators, this becomes (see App. A.11 for details)

$$H_{\text{BCS}} = \sum_{\mathbf{k}} (c_{\mathbf{k}\uparrow}^\dagger \ c_{-\mathbf{k}\downarrow}) \begin{pmatrix} \xi(\mathbf{k}) & \Delta(\mathbf{k}) \\ \Delta(\mathbf{k})^* & -\xi(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}\uparrow} \\ c_{-\mathbf{k}\downarrow}^\dagger \end{pmatrix} + \sum_{\mathbf{k}} \xi(\mathbf{k}), \quad (9.9)$$

with

$$\xi(\mathbf{k}) = -2t \cos(k) - \mu \quad \text{and} \quad \Delta(\mathbf{k}) = \Delta = \Delta_0 e^{i\delta}, \quad (9.10)$$

where $k \equiv |\mathbf{k}|$. Note that this captures the essential features of the mean-field BCS Hamiltonian Eq. (5.4) for a one-dimensional chain. If we ignore the constant energy offset $E_0 = \sum_{\mathbf{k}} \xi(\mathbf{k})$, Eq. (9.9) readily takes the reduced BdG form Eq. (9.1), where the additional indices $\alpha = \uparrow$ and $\beta = \downarrow$ label spin projections. Accordingly, Eq. (9.9) can be diagonalised as

$$\begin{aligned} H_{\text{BCS}} &= \sum_{\mathbf{k}} (c_{\mathbf{k}\uparrow}^\dagger \ c_{-\mathbf{k}\downarrow}) \begin{pmatrix} \xi(\mathbf{k}) & \Delta \\ \Delta^* & -\xi(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}\uparrow} \\ c_{-\mathbf{k}\downarrow}^\dagger \end{pmatrix} \\ &= \sum_{\mathbf{k}} (c_{\mathbf{k}\uparrow}^\dagger \ c_{-\mathbf{k}\downarrow}) \begin{pmatrix} u(\mathbf{k}) & -v(\mathbf{k}) \\ v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} \begin{pmatrix} E(\mathbf{k}) & 0 \\ 0 & -E(\mathbf{k}) \end{pmatrix} \begin{pmatrix} u(\mathbf{k}) & v(\mathbf{k}) \\ -v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}\uparrow} \\ c_{-\mathbf{k}\downarrow}^\dagger \end{pmatrix} \\ &\equiv \sum_{\mathbf{k}} (b_{\mathbf{k}\uparrow}^\dagger \ b_{-\mathbf{k}\downarrow}) \begin{pmatrix} E(\mathbf{k}) & 0 \\ 0 & -E(\mathbf{k}) \end{pmatrix} \begin{pmatrix} b_{\mathbf{k}\uparrow} \\ b_{-\mathbf{k}\downarrow}^\dagger \end{pmatrix}, \end{aligned} \quad (9.11)$$

where we defined Bogoliubov quasiparticle operators

$$\begin{pmatrix} b_{\mathbf{k}\uparrow} \\ b_{-\mathbf{k}\downarrow}^\dagger \end{pmatrix} \equiv \begin{pmatrix} u(\mathbf{k}) & v(\mathbf{k}) \\ -v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}\uparrow} \\ c_{-\mathbf{k}\downarrow}^\dagger \end{pmatrix}, \quad (9.12)$$

and where $E(\mathbf{k}), u(\mathbf{k}), v(\mathbf{k})$ take the form defined in Eqs. (9.6) and (9.4), respectively. In particular, note that we get $u(\mathbf{k}) = u(-\mathbf{k})$ and $v(\mathbf{k}) = v(-\mathbf{k})$ due to the form of $\xi(\mathbf{k})$ and $\Delta(\mathbf{k})$ given in Eq. (9.10). From Eq. (9.12) we obtain the Bogoliubov quasiparticle annihilators

$$b_{\mathbf{k}\uparrow} = u(\mathbf{k})c_{\mathbf{k}\uparrow} + v(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger \quad \text{and} \quad b_{-\mathbf{k}\downarrow} = -v(\mathbf{k})c_{\mathbf{k}\uparrow}^\dagger + u(\mathbf{k})c_{-\mathbf{k}\downarrow}, \quad (9.13)$$

both of which are needed in the construction of the fermionic Fock space of the theory. It should be noted that the notation for the Bogoliubov field operators in Eq. (9.13) is largely symbolic: each operator involves both electronic creation and annihilation operators, spin-up and spin-down projections, and positive and negative quasi-momentum quantum numbers. A Bogoliubov annihilation operator with quasi-momentum $\mathbf{k}$ and spin projection α is therefore neither an actual annihilation operator of any physical particle nor is it associated with quasi-momentum $\mathbf{k}$ or spin projection α in the original sense of these quantum numbers. Still, the notation in terms of the original quantum numbers accounts for the fact that both field operators in Eq. (9.13) are needed to capture the full fermionic complexity of the theory. Indeed, we can rewrite Eq. (9.11) as

$$H_{\text{BCS}} = \sum_{\mathbf{k},\alpha} E(\mathbf{k}) b_{\mathbf{k}\alpha}^\dagger b_{\mathbf{k}\alpha} + E'_0, \quad (9.14)$$

where $E'_0 = \sum_{\mathbf{k}} E(\mathbf{k})$ is another insignificant energy offset. Following Sec. 5.3, we write the BCS ground state as a normalised product state (see App. A.11 for details)

$$|0\rangle_b^{\text{p}} = \frac{1}{\mathcal{N}} \prod_{\mathbf{k},\alpha} b_{\mathbf{k}\alpha} |0\rangle = \prod_{\mathbf{k}} (u(\mathbf{k})e^{-i\delta} + v_0(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle, \quad (9.15)$$

where $\mathcal{N} = \prod_{\mathbf{k}} v(\mathbf{k})$ is a normalisation prefactor and $|0\rangle$ denotes the electronic vacuum defined via $c_{j\alpha}|0\rangle = 0$. If we choose to understand the phase $\delta \in \mathbb{R}/2\pi\mathbb{Z} \simeq \mathbb{S}^1$ of the superconducting gap Δ as a parameter of the BCS theory, we may construct the Berry connection one-form

$$A = i \int_b^p \langle 0 | \partial_\delta | 0 \rangle_b^p d\delta = A_\delta d\delta \quad (9.16)$$

and compute the Berry phase

$$\gamma(C) = \oint_C A = i \oint_C \int_b^p \langle 0 | \partial_\delta | 0 \rangle_b^p d\delta \quad (9.17)$$

of $|0\rangle_b$ along any closed path C in $\mathbb{S}^1$, cf. Sec. 4. To arrive at an analytical expression for this, we first determine the coefficient A_δ of the Berry connection. Using

$$\begin{aligned} \partial_\delta |0\rangle_b^p &= \partial_\delta \prod_{\mathbf{k}} (u(\mathbf{k})e^{-i\delta} + v_0(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle \\ &= -i \sum_{\mathbf{p}} u(\mathbf{p})e^{-i\delta} \prod_{\mathbf{k} \neq \mathbf{p}} \left(u(\mathbf{k})e^{-i\delta} + v_0(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) |0\rangle \end{aligned} \quad (9.18)$$

and

$$\langle 0 | \prod_{\mathbf{k} \neq \mathbf{p}} \left(u(\mathbf{k})e^{i\delta} + v_0(\mathbf{k})c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \right) \left(u(\mathbf{k})e^{-i\delta} + v_0(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) | 0 \rangle = \langle 0 | \prod_{\mathbf{k} \neq \mathbf{p}} (u(\mathbf{k})^2 + v_0(\mathbf{k})^2) | 0 \rangle = 1, \quad (9.19)$$

we get

$$\begin{aligned} A_\delta &= i \int_b^p \langle 0 | \partial_\delta | 0 \rangle_b^p \\ &\stackrel{(\diamond)}{=} \sum_{\mathbf{p}} u(\mathbf{p})e^{-i\delta} \langle 0 | \prod_{\mathbf{k}'} \left(u(\mathbf{k}')e^{i\delta} + v_0(\mathbf{k}')c_{\mathbf{k}'\uparrow} c_{-\mathbf{k}'\downarrow} \right) \prod_{\mathbf{k} \neq \mathbf{p}} \left(u(\mathbf{k})e^{-i\delta} + v_0(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) | 0 \rangle \\ &\stackrel{(*)}{=} \sum_{\mathbf{p}} u(\mathbf{p})e^{-i\delta} \langle 0 | \prod_{\mathbf{k} \neq \mathbf{p}} \left(u(\mathbf{k})e^{i\delta} + v_0(\mathbf{k})c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \right) \left(u(\mathbf{k})e^{-i\delta} + v_0(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) \left(u(\mathbf{p})e^{i\delta} + v_0(\mathbf{p})\cancel{c_{\mathbf{p}\uparrow} c_{-\mathbf{p}\downarrow}} \right) | 0 \rangle \\ &= \sum_{\mathbf{p}} u(\mathbf{p})^2 \langle 0 | \prod_{\mathbf{k} \neq \mathbf{p}} \left(u(\mathbf{k})e^{i\delta} + v_0(\mathbf{k})c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \right) \left(u(\mathbf{k})e^{-i\delta} + v_0(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) | 0 \rangle \\ &\stackrel{(*)}{=} \sum_{\mathbf{p}} u(\mathbf{p})^2, \end{aligned} \quad (9.20)$$

where we plugged in Eq. (9.18) in $(\diamond)$, used that $[c_{\mathbf{p}\uparrow} c_{-\mathbf{p}\downarrow}, c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger] = 0$ for all $\mathbf{k} \neq \mathbf{p}$ in $(*)$, and finally applied Eq. (9.19) in $(*)$. Now, inserting in Eq. (9.4) into the final line of Eq. (9.20) yields

$$\sum_{\mathbf{p}} u(\mathbf{p})^2 = \sum_{\mathbf{p}} \frac{1}{2} \left(1 + \frac{\xi(\mathbf{p})}{E(\mathbf{p})} \right) = \frac{L}{2} + \frac{1}{2} \sum_{\mathbf{p}} \frac{\xi(\mathbf{p})}{E(\mathbf{p})}, \quad (9.21)$$

where we used that the number of quasi-momenta equals the number lattice sites, i.e. $\sum_{\mathbf{p}} 1 = L$. For $\mu = 0$, we additionally have $\xi(\mathbf{k} \pm \boldsymbol{\pi}) = -\xi(\mathbf{k})$ and $E(\mathbf{k} \pm \boldsymbol{\pi}) = E(\mathbf{k})$, such that

$$\sum_{\mathbf{p}=-\boldsymbol{\pi}}^{\boldsymbol{\pi}} \frac{\xi(\mathbf{p})}{E(\mathbf{p})} = \sum_{\mathbf{p}=0}^{\boldsymbol{\pi}} \left(\frac{\xi(\mathbf{p})}{E(\mathbf{p})} + \frac{\xi(\mathbf{p}-\boldsymbol{\pi})}{E(\mathbf{p}-\boldsymbol{\pi})} \right) = \sum_{\mathbf{p}=0}^{\boldsymbol{\pi}} \left(\frac{\xi(\mathbf{p})}{E(\mathbf{p})} - \frac{\xi(\mathbf{p})}{E(\mathbf{p})} \right) = 0. \quad (9.22)$$

By definition, the closed paths C in $\mathbb{S}^1 \simeq \mathbb{R}/2\pi\mathbb{Z}$ are of the form

$$C_N : [0, 1] \rightarrow \mathbb{S}^1, \quad s \mapsto 2\pi N s \quad (9.23)$$

for some integer winding $N \in \mathbb{Z}$, cf. Sec. 2.1.8. With this, we may write the Berry phase $\gamma(C_N)$ along a closed curve C_N as

$$\gamma(C_N) = \oint_{C_N} A_\delta d\delta = \sum_{\mathbf{p}} u(\mathbf{p})^2 \left[\int_0^{2\pi N} d\delta \right] = 2\pi N \sum_{\mathbf{p}} u(\mathbf{p})^2, \quad (9.24)$$

which, using Eqs. (9.21) and (9.22), simplifies to

$$\gamma(C_N) = \pi N L \quad (9.25)$$

for $\mu = 0$.

9.3 The Sombrero Berry-Phase of the Kitaev Chain Ground State

The Kitaev chain Hamiltonian reads

$$H_{\text{Kit}} = \sum_{j=1}^{L-1} \left[-t (c_{j+1}^\dagger c_j + c_j^\dagger c_{j+1}) + \Delta c_j c_{j+1} + \Delta^* c_{j+1}^\dagger c_j^\dagger \right] - \mu \sum_{j=1}^L (c_j^\dagger c_j - \frac{1}{2}). \quad (9.26)$$

Upon Fourier transforming the electronic field operators, this becomes (see App. A.11 for details)

$$H_K = \frac{1}{2} \sum_{\mathbf{k}} (c_{\mathbf{k}}^\dagger \ c_{-\mathbf{k}}) \begin{pmatrix} \xi(\mathbf{k}) & \Delta(\mathbf{k}) \\ \Delta(\mathbf{k})^* & -\xi(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}} \\ c_{-\mathbf{k}}^\dagger \end{pmatrix} + E_0 \quad (9.27)$$

with

$$\xi(\mathbf{k}) = -2t \cos(k) - \mu \quad \text{and} \quad \Delta(\mathbf{k}) = -2i\Delta \sin(k), \quad (9.28)$$

where $k \equiv |\mathbf{k}|$. Unlike the BCS gap, the Kitaev gap is not constant, but disperses as a function of quasi-momentum. In particular, it vanishes for $k_c = 0, \pi$, which is something we must keep in mind. For future reference, we define the polar form

$$\Delta(\mathbf{k}) = -2i\Delta_0 e^{i\delta'} \sin(k) =: 2\Delta_0 \sin(k) e^{i\delta} \equiv \Delta_0(\mathbf{k}) e^{i\delta} \quad (9.29)$$

of the superconducting gap function. Here, we plugged in $\Delta = \Delta_0 e^{i\delta'}$ and absorbed the prefactor of $-i = e^{-i\frac{\pi}{2}}$, defining $\delta \equiv \delta' - \pi/2$. Note that, upon omitting the constant energy offset $E_0 = \sum_{\mathbf{k}} \xi(\mathbf{k})$, Eq. (9.27) assumes the standard BdG form, see Eq. (5.28). In particular, the absence of distinct single-particle indices $\alpha \neq \beta$ in Eq. (9.27) gives rise to the characteristic redundancy of fermionic degrees of freedom. The standard Bogoliubov diagonalisation of Eq. (9.27) then yields

$$H_K = \frac{1}{2} \sum_{\mathbf{k}} (b_{\mathbf{k}}^\dagger \ b_{-\mathbf{k}}) \begin{pmatrix} E(\mathbf{k}) & 0 \\ 0 & -E(\mathbf{k}) \end{pmatrix} \begin{pmatrix} b_{\mathbf{k}} \\ b_{-\mathbf{k}}^\dagger \end{pmatrix}, \quad (9.30)$$

where we defined Bogoliubov quasiparticle operators

$$\begin{pmatrix} b_{\mathbf{k}} \\ b_{-\mathbf{k}}^\dagger \end{pmatrix} \equiv \begin{pmatrix} u(\mathbf{k}) & v(\mathbf{k}) \\ -v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}} \\ c_{-\mathbf{k}}^\dagger \end{pmatrix}, \quad (9.31)$$

and where $E(\mathbf{k}), u(\mathbf{k}), v(\mathbf{k})$ take the form defined in Eqs. (9.6) and (9.4). Recall that in the BCS model both the $u(\mathbf{k})$ and $v(\mathbf{k})$ coefficients satisfy $u(\mathbf{k}) = u(-\mathbf{k})$ and $v(\mathbf{k}) = v(-\mathbf{k})$. Here, we find that $u(\mathbf{k}) = u(-\mathbf{k})$ but $v(\mathbf{k}) = -v(-\mathbf{k})$, because $\Delta(\mathbf{k}) \propto \sin(k)$. We can use this to show that the two Bogoliubov quasiparticle annihilation operators

$$b_{\mathbf{k}} = u(\mathbf{k})c_{\mathbf{k}} + v(\mathbf{k})c_{-\mathbf{k}}^\dagger \quad \text{and} \quad b_{-\mathbf{k}} = -v(\mathbf{k})c_{\mathbf{k}}^\dagger + u(\mathbf{k})c_{-\mathbf{k}} \quad (9.32)$$

resulting from Eq. (9.31) describe the *same* fermionic degrees of freedom: if we take the second Bogoliubov operator $b_{-\mathbf{k}}$ and invert the momentum $\mathbf{k} \mapsto -\mathbf{k}$, we simply recover the first Bogoliubov operator $b_{\mathbf{k}}$,

$$b_{-(-\mathbf{k})} = -v(\mathbf{k})c_{\mathbf{k}}^\dagger + u(\mathbf{k})c_{-\mathbf{k}} = -v(-\mathbf{k})c_{-\mathbf{k}}^\dagger + u(-\mathbf{k})c_{\mathbf{k}} = v(\mathbf{k})c_{-\mathbf{k}}^\dagger + u(\mathbf{k})c_{\mathbf{k}} = b_{\mathbf{k}}, \quad (9.33)$$

exposing their redundancy. From Eq. (9.30), we therefore get (see App. A.11 for details)

$$H_K = \sum_{\mathbf{k}} E(\mathbf{k}) b_{\mathbf{k}}^\dagger b_{\mathbf{k}} + E'_0, \quad (9.34)$$

where $E'_0 = \sum_{\mathbf{k}} E(\mathbf{k})$ is another insignificant energy offset. Following Sec. 5.3, we may write the KC ground state as a normalised product state (see App. A.11 for details)

$$|0\rangle_b^p = \prod_{\mathbf{k}} b_{\mathbf{k}} |0\rangle. \quad (9.35)$$

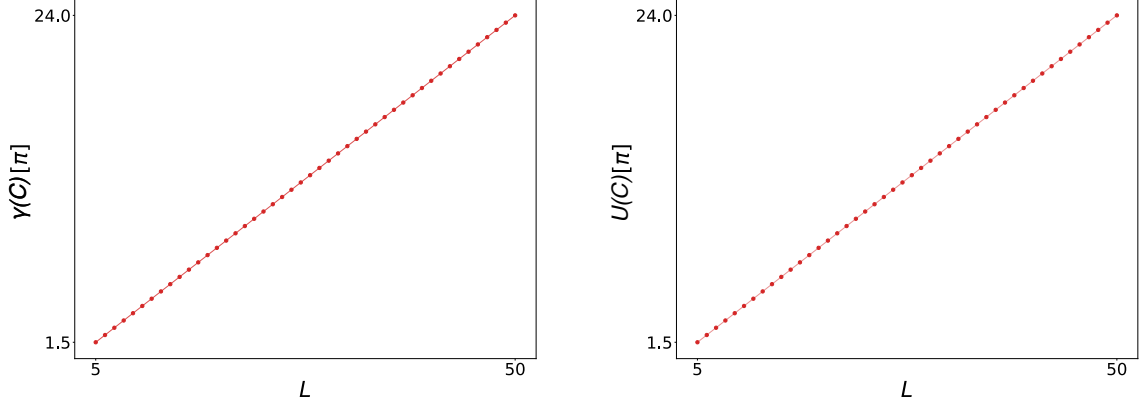


Figure 9.1: Numerical results for the Sombrero Berry phase of the Kitaev ground states as a function of system size L , with model parameters $t = \Delta_0 = 1$ and $\mu = 0$. *Left:* data obtained from the analytical expression Eq. (9.41) for $N = 1$. *Right:* data obtained using the discretised Berry phase algorithm, Eq. (4.110) with Eq. (4.112), of the BdG vacua, cf. Sec. 5.

However, recall that there exist critical quasi-momenta $\mathbf{k}_c = \mathbf{0}, \pi$, where

$$\Delta(\mathbf{k}_c) = 2i\Delta \sin(k_c) = 0, \quad (9.36)$$

so that

$$E(\mathbf{k}_c) = \xi(\mathbf{k}_c), \quad u(\mathbf{k}_c) = 1, \quad v(\mathbf{k}_c) = 0. \quad (9.37)$$

Accordingly, we find that

$$b_{\mathbf{k}_c} = u(\mathbf{k}_c)c_{\mathbf{k}_c} + v(\mathbf{k}_c)c_{-\mathbf{k}_c}^\dagger = c_{\mathbf{k}_c}. \quad (9.38)$$

i.e. that the Bogoliubov annihilation operators coincide with the electronic annihilation operators at the critical quasi-momenta. If these critical operators were included in the naive product state definition of the ground state, Eq. (9.35), they would annihilate the electronic vacuum. To avoid this, the critical quasi-momenta must be excluded from the construction, and we find (see App. A.11 for details)

$$|0\rangle_b^p = \frac{1}{\mathcal{N}} \prod_{\mathbf{k} \neq \mathbf{0}, \pm\pi} b_{\mathbf{k}} |0\rangle = \frac{1}{\mathcal{N}} \prod_{\mathbf{0} < \mathbf{k} < \pi} b_{\mathbf{k}} b_{-\mathbf{k}} |0\rangle = \prod_{\mathbf{0} < \mathbf{k} < \pi} (u(\mathbf{k})e^{-i\delta} + v_0(\mathbf{k})c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle, \quad (9.39)$$

where $\mathcal{N} = \prod_{\mathbf{0} < \mathbf{k} < \pi} v(\mathbf{k})$ is a normalisation prefactor, and $|0\rangle$ denotes the electronic vacuum once more. Note that we exclude the critical momenta $\mathbf{k}_c = \mathbf{0}, \pi$ and restrict the product over the (punctured) Brillouin zone to its positive half by pairing each $b_{\mathbf{k}}$ with its negative-momentum counterpart $b_{-\mathbf{k}}$. Apart from the restriction of the product to half of the punctured Brillouin zone, Eq. (9.39) reproduces the form of Eq. (9.15) exactly. Consequently, the expressions for the Kitaev Berry connection and Berry phase take the same form as Eqs. (9.20) and (9.24), except that all quasi-momentum sums are restricted to the positive half of the punctured Brillouin zone, i.e. we get

$$A = A_\delta d\delta \quad \text{with} \quad A_\delta = \sum_{\mathbf{0} < \mathbf{p} < \pi} u(\mathbf{p})^2 \quad (9.40)$$

and

$$\gamma(C_N) = 2\pi N \sum_{\mathbf{0} < \mathbf{p} < \pi} u(\mathbf{p})^2. \quad (9.41)$$

Using Eqs. (9.21) and (9.22), the latter simplifies to

$$\gamma(C_N) = \pi N L / 2 \quad (9.42)$$

for $\mu = 0$. The analytical expressions for the Sombrero Berry phases of the BCS and Kitaev chain ground states provide a benchmark for testing the numerical algorithm for the discretised WZ and Berry phases of Bogoliubov vacua introduced in Sec. (4.6). Figure 9.1 shows that the discretised Berry phase of the Kitaev ground states, Eq. (4.110) with Eq. (4.112), are in perfect agreement with the analytical prediction from Eq. (9.41).

10 – Exchangeless Braiding with Non-Degenerate Anyons

Topological quantum computation (TQC) is one of the most promising frameworks for fault-tolerant quantum computation (QC) [101, 114, 142, 194–203]. It is based on the exotic statistics of non-Abelian anyons, which provides topological protection against local perturbations and control errors [95, 101, 194, 204, 205]. Concretely, certain quantum systems support anyonic quasiparticles that give rise to a degenerate low-energy subspace $\mathcal{H}_0$ suitable for encoding quantum information [114]. Due to their anyonic statistics, the adiabatic exchange of these quasiparticles can be used to induce unitary transformations on $\mathcal{H}_0$. The resulting transformations are *topological* in that they depend only on the topological properties of the braids formed by the quasiparticle world lines during the exchange. Sequences of adiabatic anyon exchanges can then be applied to implement robust quantum gates on any qubit state $|\psi\rangle$ encoded in $\mathcal{H}_0$. As a result, TQC with anyons integrates error correction on a hardware level, making it a compelling architecture for robust QC applications.

In recent decades, TQC has attracted growing interest. This rise in attention is partly due to the discovery of anyonic defect modes in certain topological superconductors [14, 106, 108, 206–209]. These defect modes exhibit a range of distinctive characteristics. First, their existence is guaranteed by the bulk-defect correspondence of the underlying topological superconductor. This correspondence protects them against symmetry-preserving perturbations and pins them at zero energy [14]. As a result, the many-body ground-state energy becomes degenerate, introducing the degenerate subspace $\mathcal{H}_0$ of ground states essential for QC. Second, the defect-bound zero-energy quasiparticle excitations of the superconducting state can be understood as self-adjoint quasiparticle modes, i.e. Majorana-type fermion modes. For this reason, they are usually referred to as Majorana zero modes (MZMs). Finally, and most crucially for TQC, they are equivalent to a specific type of non-Abelian anyons known as Ising anyons [95, 101, 114, 195]. Even though Ising anyons are not universal for quantum computing – not every unitary transformation on $\mathcal{H}_0$ can be approximated to arbitrary precision by braiding transformations – their experimental accessibility and low braiding complexity¹ make the study of Ising anyons a valuable endeavour.

Proposals for concrete physical realisations include MZMs bound to vortex cores in two-dimensional $p_x + ip_y$ superconductors [106, 207, 210, 211] and to the boundaries of one-dimensional topological p -wave superconductors [108, 206, 212–222]. Further progress towards practical TQC faces a number of experimental and theoretical challenges. Experimental efforts aim to refine engineering techniques for topological superconductors and achieve the microscopic control required to prepare and manipulate MZMs, for a review see Refs. [223–231]. On the theory side, simulations must account for real-world limitations to provide more realistic expectations for experiments. Advances have been made through various methods and methodological approaches. Exact diagonalisation is effective but limited to small systems [217, 232, 233]. Studies focusing on single-particle states [216, 221] or low-energy theories [234–237] can achieve larger system sizes, but at the cost of neglecting contributions of extended bulk-states to the many-body ground state. Other efforts concentrate on time-evolving quasiparticles [117, 211, 238], applications of the Onishi determinant formula [239], and covariance matrix techniques [240–242].

Here, we consider many-body simulations for weakly coupled, finite-size Kitaev chains and propose an exchangeless braiding protocol, in which the unitary braiding transformation is driven not by an adiabatic exchange of MZMs, but by 2π rotations of the superconducting phase ϕ that keep them stationary in real space [200, 243–246]. This was originally proposed by Kitaev [206] and later picked up by others, including Fu, Teo and Kane [247, 248], Chiu *et al.* [14], and Sanno *et al.* [211]. The Kitaev chain model [206] has recently attracted renewed attention because it was found to provide a suitable low-energy description of some semi-conductor/superconductor and spin-chain/superconductor hybrid systems [210, 213].

¹In the end of Sec. 6.3.2, specifically in Eqs. (6.66) and (6.67), we saw that a simple double exchange of Ising anyons is enough to implement X - and Z -gates. This is what we mean by low braiding complexity – simple gates can be achieved through simple braiding protocols. Compare, in particular, to the high braiding complexity of (the *universal* class of) Fibonacci anyons that we mentioned in the final paragraph of Sec. 6.3.1.

Local manipulations of the superconducting phase are generally difficult to carry out. Here, their use is motivated by Yu-Shiba-Rusinov (YSR) realisations of topological superconductivity [213, 249–253]. In such setups, a ferromagnetically aligned classical spin array is exchange-coupled to a conventional s -wave superconductor, creating low-energy YSR bands that reside inside the host’s superconducting gap and inherit a *topological* p -wave superconducting order by proximity to the s -wave host [213, 249]. A collective rotation of the classical spin array then translates into a local phase rotation of the proximity-induced p -wave order parameter in the YSR bands, as suggested in Ref. [213].

In order to detect Ising-anyon type MZM braiding, we analyse the non-Abelian Wilczek–Zee (WZ) phase matrix [254] of the low-energy subspace $\mathcal{H}_0(\phi)$ under cyclic evolutions $\phi \mapsto \phi + 2\pi$ of the superconducting phase. That is, we explicitly evaluate a $U(n)$ holonomy associated to the n -dimensional space $\mathcal{H}_0(\phi)$, focusing solely on the geometric phase and neglecting dynamical contributions. Methodologically, we follow Ref. [94] and use the Bertsch–Robledo–Pfaffian overlap formula [87, 90–93] to circumvent the Onishi sign problem in the computation of many-body overlaps between zero-energy (low-energy) Fock states in the numerical computation of WZ phases [94, 210].

Based on a numerical analysis of the WZ phase, we demonstrate that the MZMs of the Kitaev chain retain their anyonic properties even in the presence of finite-size induced MZM-interactions that lift their degeneracy significantly. This stability enables the implementation of an exchangeless double-braiding protocol driven by cyclic rotations of the superconducting phase. Our main focus is on composite systems of two weakly linked Kitaev chains. There, we find that the relative strength of couplings between MZMs within the same subchain and between different subchains controls whether an exchangeless double braiding within either subchain produces a unitary Z - or X -gate transformation.

The remainder of this chapter is organised as follows. In Secs. 10.1 and 10.2, we introduce the model and describe the technical framework behind our numerical approach. Subsequently, Sec. 10.3 presents a numerical analysis of the Wilczek–Zee phase to demonstrate exchangeless double braiding and quantify finite-size effects in a single Kitaev chain. In the final section, Sec. 10.4, we extend this analysis to a simple two-Kitaev chain network, where the exchangeless double braiding process is shown to induce Z - or X -gate transformations, depending on the parameters of the model.

Throughout this chapter, we closely follow our original presentation in [RQ3].

10.1 The Kitaev Chain Model

The Kitaev chain is a minimal model for a topological superconductor. Proposed by Alexei Kitaev in 2001 [206], it describes a one-dimensional lattice of itinerant spinless fermions with a superconducting p -wave pairing. In the topologically non-trivial phase, it features unpaired Majorana zero modes (MZMs) at its boundary. The second-quantised Kitaev chain Hamiltonian reads

$$H_{\text{Kit}} = \sum_{j=1}^{L-1} \left[-t (c_{j+1}^\dagger c_j + c_j^\dagger c_{j+1}) + \Delta c_j c_{j+1} + \Delta^* c_{j+1}^\dagger c_j^\dagger \right] - \mu \sum_{j=1}^L (c_j^\dagger c_j - \frac{1}{2}), \quad (10.1)$$

where the index j labels the sites of a one-dimensional chain of length L . The first term of Eq. (10.1) is the generic tight-binding hopping. It is governed by the NN hopping strength t , which we set to $t \equiv 1$ to fix the energy unit. The third term in Eq. (10.1) implements a chemical potential of strength μ and can be used to tune between topologically distinct phases. The remaining terms represent the superconducting p -wave pairing determined by a complex gap parameter

$$\Delta = \Delta_0 e^{i\phi} = \Delta_0 (\cos(\phi) + i \sin(\phi)), \quad (10.2)$$

where Δ_0 and ϕ denote the real amplitude and phase, respectively. The pairing term explicitly breaks the $U(1)$ charge symmetry, such that the total number $N = \sum_j c_j^\dagger c_j$ of spinless fermions is no longer conserved. Fermion parity $(-1)^N$ remains a good quantum number.

The beauty of the Kitaev chain is that it provides a simple, instructive argument for the existence of unpaired boundary MZMs. To see this, we consider the special case with $\phi = 0$, so that $\Delta = \Delta_0$. Following Sec. 5.5, we may introduce self-adjoint Majorana operators

$$\gamma_j^A = c_j^\dagger + c_j = \gamma_j^{A\dagger} \quad \text{and} \quad \gamma_j^B = i(c_j^\dagger - c_j) = \gamma_j^{B\dagger}, \quad (10.3)$$

which allow us to rewrite Eq. (10.1) with $\Delta = \Delta_0$ as (for details see App. A.12)

$$H_{\text{Kit}} = \frac{i}{2} \left(\sum_{j=1}^{L-1} [(\Delta_0 + t) \gamma_j^B \gamma_{j+1}^A + (\Delta_0 - t) \gamma_j^A \gamma_{j+1}^B] - \mu \sum_{j=1}^L \gamma_j^A \gamma_j^B \right). \quad (10.4)$$

Based on this *Majorana representation* of H_{Kit} we may identify two special phases of the model. In the first phase (A) we have $\Delta_0 = t = 0$ and $\mu \neq 0$. This yields

$$H_{\text{Kit}} = -\frac{i\mu}{2} \sum_{j=1}^L \gamma_j^A \gamma_j^B, \quad (10.5)$$

which, by Eq. (10.3), is equivalent to

$$H_{\text{Kit}} = -\mu \sum_{j=1}^L \left(c_j^\dagger c_j - \frac{1}{2} \right). \quad (10.6)$$

Note that Eq. (10.5) only couples Majorana operators γ_j^A and γ_j^B from the *same* lattice site j to one another. If $\mu < 0$, Eq. (10.6) reveals that ground state $|\text{GS}\rangle$ must satisfy $c_j |\text{GS}\rangle = 0$ for all $j = 1, \dots, L$, i.e. it is simply given by the trivial fermion vacuum

$$|\text{GS}\rangle = |0\rangle. \quad (10.7)$$

For $\mu > 0$, an analogous argument applies to the fermionic anti-vacuum $|\bar{0}\rangle$ defined via $c_j^\dagger |\bar{0}\rangle = 0$ for all $j = 1, \dots, L$. The second phase (B) turns out much more interesting. It is distinguished by $\Delta_0 = t > 0$ and $\mu = 0$, with which Eq. (10.4) becomes

$$H_{\text{Kit}} = it \sum_{j=1}^{L-1} \gamma_j^B \gamma_{j+1}^A. \quad (10.8)$$

This configuration of the model only couples Majorana operators γ_j^B and γ_{j+1}^A from *different* lattice sites. In particular, the Majorana operators γ_1^A and γ_L^B remain *unpaired* – they do not enter the Hamiltonian in Eq. (10.8) at all. This means that they constitute self-adjoint zero energy modes, i.e. MZMs, that are perfectly localised on the chain boundary sites $j = 1$ and $j = L$. As a consequence, we can combine γ_1^A and γ_L^B into a single non-local complex fermion

$$b_0 := \frac{1}{2} (\gamma_1^A + i\gamma_L^B). \quad (10.9)$$

which is perfectly localised on the opposite ends of the chain and has strictly zero energy. The latter is evident from the fact that b_0 does not appear in the Hamiltonian

$$H_{\text{Kit}} = 2t \sum_{j=1}^{L-1} \left(b_j^\dagger b_j - \frac{1}{2} \right), \quad (10.10)$$

that we obtain from Eq. (10.8) upon defining

$$b_j = \frac{1}{2} (\gamma_{j+1}^A + i\gamma_j^B) \quad \text{and} \quad b_j^\dagger = \frac{1}{2} (\gamma_{j+1}^A - i\gamma_j^B). \quad (10.11)$$

If $t > 0$, the ground state $|\text{GS}\rangle$ of Eq. (10.10) is given by the b -vacuum,

$$|\text{GS}\rangle = |0\rangle_b, \quad (10.12)$$

defined by $b_j |0\rangle_b = 0$ for all $j = 1, \dots, L-1$. However, the existence of the unpaired mode b_0 means that there is a second ground state, namely $b_0^\dagger |0\rangle_b$, which has the same energy as $|\text{GS}\rangle$ and causes all states of the system to be doubly degenerate. Note that the complex fermion associated with b_0 is a peculiar object: it is perfectly localised at the two boundary sites $j = 1$ and $j = L$, but at the same time extremely non-local if the chain length L becomes large. We are soon going to identify the two boundary MZMs of phase (B) as the topological edge modes of the model.

The Kitaev chain Hamiltonian from Eq. (10.1) constitutes a BdG Hamiltonian,

$$H_{\text{Kit}} = \frac{1}{2} \Psi^\dagger h_{\text{Kit}} \Psi, \quad (10.13)$$

where $\Psi = (\mathbf{c} \mathbf{c}^\dagger)^\top$ denotes the Nambu spinor of fermionic annihilation and creation operators, and

$$h_{\text{Kit}} = \begin{pmatrix} T & \Delta^\dagger \\ \Delta & -T^* \end{pmatrix} \quad (10.14)$$

is given in terms of the single-particle hopping matrix T and the superconducting gap matrix Δ with elements

$$T_{jk} = -t(\delta_{(j+1)k} + \delta_{j(k+1)}) - \mu\delta_{jk} \quad \text{and} \quad \Delta_{jk} = \Delta(\delta_{(j+1)k} - \delta_{j(k+1)}). \quad (10.15)$$

As such, H_{Kit} possesses the tautological particle-hole structure

$$\bar{\Xi} H_{\text{Kit}} \bar{\Xi}^\dagger = -H_{\text{Kit}}, \quad (10.16)$$

that was discussed around Eq. (5.50) in Sec. 5.2. Here, $\bar{\Xi}$ with $\bar{\Xi}^2 = +1$ is the anti-unitary tautological particle-hole conjugation operator from Eq. (3.36). Beyond this, the Kitaev chain Hamiltonian also exhibits a genuine particle-hole symmetry²

$$\mathcal{C} H_{\text{Kit}} \mathcal{C}^\dagger = H_{\text{Kit}}, \quad (10.17)$$

implemented by the anti-unitary PHS operator $\mathcal{C}$ with $\mathcal{C}^2 = +1$, that was defined in Eq. (3.42). It transforms the elementary fermionic annihilation and creation operators as

$$\mathcal{C} c_j \mathcal{C}^\dagger = (-1)^j c_j^\dagger, \quad \mathcal{C} i \mathcal{C}^\dagger = -i. \quad (10.18)$$

As we have seen in Sec. (9.3), the Kitaev chain Hamiltonian H_{Kit} can be diagonalised in $\mathbf{k}$ -space by means of a Bogoliubov transformation. Concretely, the Fourier transform $c_j = 1/\sqrt{L} \sum_{\mathbf{k}} e^{-i\mathbf{k}\mathbf{R}_j} c_{\mathbf{k}}$ of the elementary field operators allows us to write H_{Kit} as

$$H_{\text{Kit}} = \sum_{\mathbf{k}} \phi^\dagger(\mathbf{k}) h_{\text{Kit}}(\mathbf{k}) \phi(\mathbf{k}), \quad (10.19)$$

where we defined the Nambu spinor $\phi(\mathbf{k}) = (c_{\mathbf{k}} c_{-\mathbf{k}}^\dagger)^\top$ of Fourier transformed annihilation and creation operators with quasi-momenta $\mathbf{k}$ and $-\mathbf{k}$, respectively. The 2×2 Bloch matrix in Eq. (10.19) takes the form

$$h_{\text{Kit}}(\mathbf{k}) = \mathbf{h}(\mathbf{k}) \boldsymbol{\sigma}, \quad (10.20)$$

where $\boldsymbol{\sigma}$ denotes the vector of Pauli matrices σ_j acting on particle and hole components of the Nambu spinors. The coefficient functions $\mathbf{h}(\mathbf{k}) = (h_x(\mathbf{k}) h_y(\mathbf{k}) h_z(\mathbf{k}))$ are given by (for details see App. A.12)

$$h_x(\mathbf{k}) = 2\Delta_0 \sin(\phi) \sin(k), \quad h_y(\mathbf{k}) = -2\Delta_0 \cos(\phi) \sin(k), \quad h_z(\mathbf{k}) = \mu + 2t \cos(k), \quad (10.21)$$

where $k \equiv |\mathbf{k}|$. The Bogoliubov diagonalisation

$$H_K = \frac{1}{2} \sum_{\mathbf{k}} (b_{\mathbf{k}}^\dagger \ b_{-\mathbf{k}}) \begin{pmatrix} E(\mathbf{k}) & 0 \\ 0 & -E(\mathbf{k}) \end{pmatrix} \begin{pmatrix} b_{\mathbf{k}} \\ b_{-\mathbf{k}}^\dagger \end{pmatrix} \quad (10.22)$$

²The Kitaev chain is often said to possess a TRS $\mathcal{T}$ with $\mathcal{T}^2 = +1$, acting as $\mathcal{T} c_j \mathcal{T}^\dagger = c_j$ and $\mathcal{T} i \mathcal{T}^\dagger = -i$. While this interpretation mathematically consistent, it clashes with the physical interpretation of the c_j as real fermions: physical fermions carry spin, leading to $\mathcal{T}^2 = (-1)^N$, so for single-particle physics ($N = 1$) we expect $\mathcal{T}^2 = -1$. In $\text{SU}(2)$ symmetric systems, one may recover single-particle TRS with $\mathcal{T}^2 = +1$ after an $\text{SU}(2)$ symmetry reduction. However, this is not possible in spinless models like the Kitaev chain. One must therefore either (i) interpret the c_j as spinless “toy” fermions and retain the TRS perspective, or (ii) view them as spin-polarised fermions in an $\text{SU}(2)$ -broken setting and reinterpret the symmetry properties in terms of a PHS instead. Details on this are provided in Ref. [78].

of Eq. (10.19) yields Bogoliubov quasiparticle operators

$$\begin{pmatrix} b_{\mathbf{k}} \\ b_{-\mathbf{k}}^\dagger \end{pmatrix} \equiv \begin{pmatrix} u(\mathbf{k}) & v(\mathbf{k}) \\ -v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}} \\ c_{-\mathbf{k}}^\dagger \end{pmatrix}, \quad (10.23)$$

where $u(\mathbf{k}), v(\mathbf{k})$ take the form defined in Eq. (9.4), and energy bands

$$E_{\pm}(\mathbf{k}) = \pm \sqrt{h_x(\mathbf{k})^2 + h_y(\mathbf{k})^2 + h_z(\mathbf{k})^2} = \pm |\mathbf{h}(\mathbf{k})|, \quad (10.24)$$

separated by an energy gap

$$\Delta E := \min_{\mathbf{k}} (E_+(\mathbf{k}) - E_-(\mathbf{k})) = 2 \cdot \min_{\mathbf{k}} |\mathbf{h}(\mathbf{k})|. \quad (10.25)$$

In contrast to the previously discussed honeycomb models, where the lattice symmetry guarantees that the smallest band gap appears at the high symmetry Dirac points, the quasimomentum minimising Eq. (10.25) depends on the model parameters. Specifically, we find (for details see App. A.12) energy gaps

$$\Delta E = 2 \cdot \min \left(|\mu + 2t|, |\mu - 2t|, |\Delta_0| \sqrt{4 + \frac{\mu^2}{(\Delta_0^2 - t^2)}} \right). \quad (10.26)$$

with

$$\Delta E = \begin{cases} 2|\mu + 2t| & \text{attained at } k = 0 \\ 2|\mu - 2t| & \text{attained at } k = \pm\pi \\ 2|\Delta_0| \sqrt{4 + \frac{\mu^2}{(\Delta_0^2 - t^2)}} & \text{attained at } k = \pm \arccos \left(\frac{t\mu}{2(\Delta_0^2 - t^2)} \right), \end{cases} \quad (10.27)$$

respectively. Of course, the solution at $k = \pm \arccos \left(\frac{t\mu}{2(\Delta_0^2 - t^2)} \right)$ only exists if

$$\left| \frac{t\mu}{2(\Delta_0^2 - t^2)} \right| \leq 1. \quad (10.28)$$

For $\Delta_0 > 0$, the energy gap vanishes at $k = 0$ ($k = \pm\pi$) when $\mu = -2t$ ($\mu = 2t$). If $\Delta_0 = 0$, we additionally get gap closures at

$$k = \pm \arccos \left(-\frac{\mu}{2t} \right), \quad (10.29)$$

whenever $|\mu| \leq 2|t|$. As before, the nodal surfaces defined by $|\mu| = 2|t|$ divide the three-dimensional parameter space spanned by t, μ and Δ_0 into separate regions with finite band gaps $\Delta E > 0$. These regions are again closely related to the topologically trivial and non-trivial phases of the Kitaev model. Since the NN hopping strength t usually sets the energy scale, the onsite potential μ is often used to tune between the topologically distinct phases.

Following Sec. 9.3, we may write the Kitaev chain ground state $|\text{GS}\rangle$ as a normalised product state (see App. A.11 for details)

$$|\text{GS}\rangle = \frac{1}{\mathcal{N}} \prod_{0 < \mathbf{k} < \pi} b_{\mathbf{k}} |0\rangle \equiv \frac{1}{\mathcal{N}} \bigwedge_{0 < \mathbf{k} < \pi} |b(\mathbf{k})\rangle, \quad (10.30)$$

in which all positive-energy quasiparticle modes associated with $b_{\mathbf{k}}$ are empty. Here, $\mathcal{N} = \prod_{0 < \mathbf{k} < \pi} v(\mathbf{k})$ is a normalisation prefactor, $|0\rangle$ denotes the fermionic reference vacuum defined by $c_{\mathbf{k}} |0\rangle = 0$ for all $\mathbf{k} \in \mathbb{T}_{\mathbf{k}}^1$, and $|b(\mathbf{k})\rangle \equiv b_{\mathbf{k}} |0\rangle$. Note that the form of the BdG ground state in Eq. (10.30) looks very similar to that of the tight-binding ground states given in Eqs. (7.24) and (8.13). In fact, we may reinterpret Eq. (10.30) as the state

$$|\text{GS}\rangle = \frac{1}{\mathcal{N}} \prod_{0 < \mathbf{k} < \pi} d_{\mathbf{k}}^\dagger |0\rangle \equiv \frac{1}{\mathcal{N}} \bigwedge_{0 < \mathbf{k} < \pi} |d^\dagger(\mathbf{k})\rangle, \quad (10.31)$$

in which all *negative-energy* quasi-hole modes associated with $b_{\mathbf{k}}^\dagger \equiv d_{\mathbf{k}}$ are *occupied*. This representation directly mirrors the familiar Slater determinant structure of Eqs. (7.24) and (8.13), in which the valence single-particle modes are filled. Recall, however, that the physical meaning here is quite different: the $b_{\mathbf{k}}$ are coherent superpositions of creation and annihilation operators, so the BdG vacuum is not a filled valence band in the usual sense, but rather the unique state of Cooper pairs annihilated by all positive-energy quasiparticles, or equivalently, filled with all negative-energy ones.

Despite this distinction, the form of Eq. (10.30) still makes it clear that the topology of the Bogoliubov quasiparticle states $b_{\mathbf{k}}$ encodes the topological properties of the many-body ground state.

10.1.1 Topology of the Kitaev Chain

The Bogoliubov quasiparticle states $|b(\mathbf{k})\rangle$ form a complex line bundle $\mathcal{B} \xrightarrow{\pi} \mathbb{S}_{\mathbf{k}}^1$ over the one-dimensional Brillouin torus $\mathbb{S}_{\mathbf{k}}^1 = \mathbb{T}_{\mathbf{k}}^1$. Typically, we would base the topological characterisation of a complex vector bundle $E \xrightarrow{\pi} B$ on its universal Chern classes $c_i \in H^{2i}(B, \mathbb{Z})$, cf. Sec. 2.3.3. However, in the present case, the base manifold $B = \mathbb{S}_{\mathbf{k}}^1$ is *odd-dimensional*, which limits the utility of the Chern classes.

Generally, Chern numbers – cf. Sec. 2.3.6 – cannot be defined for complex vector bundles over odd-dimensional base manifolds. The reason for this is simply that Chern classes live in even-degree cohomology groups, while the orientation class of an odd-dimensional base manifold belongs to an odd-degree homology group. For $\dim_{\mathbb{R}}(B) = 1$, the situation is even more restrictive: the total Chern class $c(\mathcal{B})$ becomes trivial, $c(\mathcal{B}) = 1$, since all of its components $c_i(\mathcal{B})$ vanish for $i > 1$, cf. Def. 2.3.4.

This difficulty motivates the consideration of another type of characteristic classes, known as secondary characteristic classes. These are called so because they capture topological information precisely in situations where the primary classes, such as the Chern classes, become trivial. To see this, we adopt the perspective of Chern–Weil theory, where characteristic classes are represented in de Rham cohomology by invariant polynomials applied to the curvature $\mathcal{F}$ of a connection $\mathcal{A}$, cf. Sec. 2.3.5. Consider an arbitrary characteristic class $x_j(\mathcal{F}) \in H^{2j}(B, \mathbb{R})$ defined as the degree- j component of a characteristic polynomial $x(\mathcal{F})$. Since $x_j(\mathcal{F})$ is closed, it can *locally*, i.e. on a chart $U_k \subset B$, be written as an exact form

$$x_j(\mathcal{F})|_{U_k} = dQ_{2j-1}^x(\mathcal{A}_k), \quad (10.32)$$

where $Q_{2j-1}^x(\mathcal{A}_k) \in \Gamma(T^*U_k \otimes \mathfrak{g}, U_k)$ is a Lie-algebra valued $(2j-1)$ -form known as the Chern–Simons form of $x_j(\mathcal{F})$ [39]. If the underlying “primary” characteristic class $x_j(\mathcal{F})$ vanishes,³ Eq. (10.32) shows that $Q_{2j-1}^x(\mathcal{A}_k)$ *itself* becomes closed. In this case, it defines a *new* cohomology class in one degree lower, $Q_{2j-1}^x(\mathcal{A}) \in H^{2j-1}(B, \mathbb{R})$, which is often called Chern–Simons secondary characteristic class associated with $x_j(\mathcal{F})$. This characteristic class can then be paired against the base manifold B to give a secondary characteristic number

$$\text{CS}_{2j-1}^x[\mathcal{A}] \equiv \langle [Q_{2j-1}^x(\mathcal{A})], [B] \rangle = \oint_B Q_{2j-1}^x(\mathcal{A}), \quad (10.33)$$

that is referred to as the Chern–Simons invariant $\text{CS}_{2j-1}^x[\mathcal{A}]$ of $x(\mathcal{F})$. Consider, for example, the Chern–Simons form of the j -th Chern character $ch_j(\mathcal{F}) \in H^{2j}(B, \mathbb{Q})$. It reads [14, 39]

$$Q_{2j-1}^{ch}(\mathcal{A}) = \frac{1}{(j-1)!} \left(\frac{i}{2\pi} \right)^j \int_0^1 dt \, \text{tr}(\mathcal{A} \mathcal{F}_t^{j-1}), \quad (10.34)$$

where $\mathcal{F}_t := t\mathcal{F} + (t^2 - t)\mathcal{A}^2$. In the first few cases $j = 1$ and $j = 2$, this gives

$$Q_1^{ch}(\mathcal{A}) = \frac{i}{2\pi} \text{tr}(\mathcal{A}) \quad \text{and} \quad Q_3^{ch}(\mathcal{A}) = \frac{1}{2} \left(\frac{i}{2\pi} \right)^2 \text{tr} \left(\mathcal{A}_k d\mathcal{A} + \frac{2}{3} \mathcal{A}^3 \right). \quad (10.35)$$

In the tenfold classification of topological quantum matter with symmetries [14], the ground state of the Kitaev chain is characterised by the Chern–Simons invariant

$$\text{CS}_1^{ch}[\mathcal{A}] = \frac{i}{2\pi} \int_{-\pi}^{\pi} \text{tr}(\mathcal{A}(\mathbf{k})) \quad (10.36)$$

³Reasons for vanishing primary characteristic classes include (i) a vanishing curvature $\mathcal{F} = 0$, say due to symmetry constraints like TRS, and (ii) exceeding dimensional restrictions, like $c_i(E) = 0$ if $i > \dim_{\mathbb{C}}(F)$ or $2i > \dim_{\mathbb{R}}(B)$ for Chern classes, cf. Def. 2.3.4.

of the first Chern character $ch_1(\mathcal{F})$. Here, $\mathcal{A}(\mathbf{k}) = \langle b(\mathbf{k}) | \partial_{\mathbf{k}} b(\mathbf{k}) \rangle$ denotes the Berry connection of the Bogoliubov quasiparticle states from Eq. (10.30). This gives rise to the $\mathbb{Z}_2$ Kitaev chain invariant [14]

$$\nu = \exp \left[2\pi i \text{CS}_1^{\text{ch}}[\mathcal{A}] \right] = \begin{cases} +1 & \text{for } |\mu| > 2|t| \quad (\text{trivial}) \\ -1 & \text{for } |\mu| < 2|t| \quad (\text{non-trivial}) \end{cases} \quad (10.37)$$

There exists a powerful alternative formulation of this invariant. Namely, it can be expressed as the “degree” of the map (for details see App. A.12)

$$m : \mathbb{S}_{\mathbf{k}}^1 \rightarrow \mathbb{S}_m^1 \subset \mathbb{R}^3, \quad \mathbf{k} \mapsto \frac{\mathbf{h}(\mathbf{k})}{|\mathbf{h}(\mathbf{k})|}, \quad (10.38)$$

where $\mathbf{h}(\mathbf{k})$ is the vector of coefficient functions Eq. (10.21) defined in Eq. (10.20). Note that Eq. (10.38) is well-defined if and only if $|\mathbf{h}(\mathbf{k})| > 0$ for all $\mathbf{k}$. Since the band gap Eq. (10.25) is directly determined by $|\mathbf{h}(\mathbf{k})|$, Eq. (10.38) is well-defined if and only if the system is gapped, reflecting the conventional understanding of topological phases of matter. The existence of $m(\mathbf{k})$ is therefore directly tied to the presence of a bulk energy gap.

The degree of a continuous map like $m : \mathbb{S}_{\mathbf{k}}^1 \rightarrow \mathbb{S}_m^1$ is an integer that counts, roughly speaking, how many times the domain manifold winds around the target manifold. More generally, consider two closed, connected, orientable n -manifolds X and Y . Any continuous map $f : X \rightarrow Y$ induces a homomorphism $f_* : H_n(X) \rightarrow H_n(Y)$ between their top homology groups. By orientability of X and Y , we have $H_n(X) \simeq H_n(Y) \simeq \mathbb{Z}$, so that f_* is a group homomorphism $f_* : \mathbb{Z} \rightarrow \mathbb{Z}$. Thus, it must be of the form

$$f_*(x) = \alpha x, \quad (10.39)$$

where $\alpha \in \mathbb{Z}$ is some fixed integer. This integer is a homotopy invariant known as the degree $\deg f \equiv \alpha$ of $f : X \rightarrow Y$. For continuous maps from the n -sphere to itself, like the map $m(\mathbf{k})$ defined in Eq. (10.38), the degree is a *complete* homotopy invariant, meaning that two maps $f, g : \mathbb{S}^n \rightarrow \mathbb{S}^n$ are homotopic *if and only if* $\deg f = \deg g$. The completeness of the degree $\deg m$ of $m(\mathbf{k})$ as a homotopy invariant is what makes this formulation of the invariant so valuable. It tells us that two Kitaev systems are topologically equivalent *if and only if* their topological invariants are the same.

In order to determine the degree of a map in practice we turn to de Rham cohomology, which provides us with the relation

$$\deg f \cdot \int_Y \omega = \int_X f^* \omega, \quad (10.40)$$

where ω is some n -form on Y and $f^* \omega$ is an n -form on X called the pullback of ω to by f . Being interested in the value of $\deg f$, it is natural to compute this for a normalised volume form η with $\int_Y \eta = 1$, since this immediately yields

$$\deg f = \int_X f^* \eta. \quad (10.41)$$

If we use the above formula to compute the degree of $m : \mathbb{S}_{\mathbf{k}}^1 \rightarrow \mathbb{S}_m^1$ from Eq. (10.38), we find (for details see App. A.12)

$$|\deg m| = \begin{cases} 0 & \text{for } \Delta_0 = t = 0, \mu \neq 0 \\ 1 & \text{for } \Delta_0 = t \neq 0, \mu = 0 \end{cases} \quad (10.42)$$

in the two prototypical phases (A) and (B) that we discussed in Eqs. (10.5) and (10.8), respectively. Since the degree remains invariant under continuous parameter changes that keep the energy gap open, the $\mathbb{Z}_2$ Chern–Simons invariant from Eq. (10.37) can be expressed as

$$\nu = (-1)^{|\deg m|} = \begin{cases} +1 & \text{for } |\mu| > 2|t| \quad (\text{trivial}) \\ -1 & \text{for } |\mu| < 2|t| \quad (\text{non-trivial}) \end{cases} \quad (10.43)$$

We have seen in Eq. (10.8) that the topologically non-trivial Kitaev chain features zero-energy Majorana modes on its boundary. In the special parameter configuration (B), these modes are perfectly localised on the boundary sites $j = 1$ and $j = L$ of the chain. As a result, they are exactly degenerate at zero energy, regardless of the chain length L . For different parameter choices within the topological phase, the modes are only exponentially localised at the boundary sites and therefore hybridise in chains with finite lengths $L < \infty$. Specifically, the Majorana zero modes take the form [206]

$$\Gamma^A = \sum_{j=1}^L (\alpha_+ x_+^j + \alpha_- x_-^j) \gamma_j^A \quad \text{and} \quad \Gamma^B = \sum_{j=1}^L (\beta_+ x_+^j + \beta_- x_-^j) \gamma_{L+1-j}^B, \quad (10.44)$$

where $\alpha_{\pm}$ and $\beta_{\pm}$ are real coefficients and $x_{\pm}^j$ is the j -th power of the so-called “decay parameters” [206]

$$x_{\pm} = x_{\pm}(t, \mu, \Delta_0) = \frac{-\mu \pm \sqrt{\mu^2 - 4(t^2 - \Delta_0^2)}}{2(t + \Delta_0)}. \quad (10.45)$$

The interaction between Γ^A and Γ^B is described by an effective two-mode Hamiltonian

$$H_{\text{eff}} = \frac{i}{2} K \Gamma^A \Gamma^B, \quad (10.46)$$

where the interaction strength K is roughly given by

$$K \equiv K(t, \mu, \Delta_0; L) \propto \exp \left[-\frac{L}{\ell(t, \mu, \Delta_0)} \right] \quad (10.47)$$

with the characteristic length [206]

$$\ell \equiv \ell(t, \mu, \Delta_0) = \left[\min \{ |\ln(|x_+|)|, |\ln(|x_-|)| \} \right]^{-1}. \quad (10.48)$$

A motivation of Eqs. (10.44) and (10.45) is given in App. A.12. Note how the “perfect” topological parameter configuration with $\Delta_0 = t \neq 0$ and $\mu = 0$ yields decay parameters $x_{\pm} = 0$, so that $\ell = 0$ and, accordingly, $K = 0$. In fact, we also get $\Gamma^A = \Gamma^B = 0$, which accounts for the fact that the boundary MZMs at $\Delta_0 = t \neq 0$ and $\mu = 0$ are perfectly localised on the boundary sites, rather than “only” exponentially, so that Eq. (10.44) does not apply. Conversely, the “perfect” trivial parameter configuration with $\Delta_0 = t = 0$ and $\mu \neq 0$ leaves the decay parameters ill-defined instead.

In order to gain an understanding of the (presumed) boundary MZMs across the extended topologically trivial and non-trivial phases, we compare the expressions in Eqs. (10.44) and (10.45) for more general parameter sets from the trivial ($|\mu| > 2|t|$) and non-trivial ($|\mu| < 2|t|$) parameter regimes, respectively. We find:

(A) For $2|t| < |\mu|$, we either get

$$|x_-| < 1 < |x_+| \quad \text{or} \quad |x_+| < 1 < |x_-|. \quad (10.49)$$

Since Eq. (10.44) must remain normalisable in the macroscopic limit $L \rightarrow \infty$, this implies that only one pair of coefficients, (α_-, β_+) or (α_+, β_-) , can be non-zero. Consequently, the boundary conditions (for details see App. A.12)

$$\alpha_+ x_+^0 + \alpha_- x_-^0 = 0 \quad \text{and} \quad \beta_+ x_+^0 + \beta_- x_-^0 = 0, \quad (10.50)$$

which guarantee that boundary modes cannot leave the system, reduce to

$$\alpha_- x_-^0 = 0 \text{ and } \beta_+ x_+^0 = 0 \quad \text{or} \quad \alpha_+ x_+^0 = 0 \text{ and } \beta_- x_-^0 = 0. \quad (10.51)$$

The fact that these can only be satisfied by trivial coefficients, tells us that the supposed boundary zero modes cannot exist in the extended topologically trivial phase.

(B) For $2|t| > |\mu|$, on the other hand, we get

$$|x_-|, |x_+| < 1, \quad (10.52)$$

which *does* allow for non-trivial boundary conditions, Eqs. (10.50), of the coefficients (α_+, α_-) and (β_+, β_-) . This means that the boundary zero modes can, and generally will, exist in the extended topologically non-trivial phase. In particular, Γ^A is localised near site $j = 1$, while Γ^B is localised near site L .

10.2 Computational Methodology

The Kitaev chain Hamiltonian in Eq. (10.1) can be understood as a continuous family $H_{\text{Kit}}(\phi)$ of Hamiltonians parameterised by the phase $\phi \in \mathbb{S}_\phi^1$ of the superconducting gap parameter $\Delta = \Delta_0 e^{i\phi}$. From $H_{\text{Kit}}(\phi) |n(\phi)\rangle = E_n(\phi) |n(\phi)\rangle$ we also obtain a continuous family of instantaneous many-body energy eigenstates $|n(\phi)\rangle$, which form a vector bundle over the parameter manifold $\mathbb{S}_\phi^1$. In the topologically non-trivial phase, the Kitaev chain features boundary MZMs that obey (projective) non-Abelian Ising anyon statistics, cf. Sec. 6.3.2. As described in Sec. 6.4, the presence of these boundary MZMs results in a twofold degeneracy of the many-body ground state energy. Naturally, this ground-state degeneracy will be higher when networks of more than one Kitaev chain are considered. Note that we will continue to refer to the low-energy Majorana modes as MZMs even if they acquire a finite energy splitting due to hybridisation in finite-length chains or weakly coupled chain segments.

Rather than adiabatically exchanging MZMs in real space [94, 214–217, 220–222] we consider *exchangeless* braiding driven by continuous changes of the superconducting phase ϕ by integer multiples of 2π . To demonstrate this, we consider the d_0 -dimensional subspace $\mathcal{H}_0(\phi)$ spanned by the instantaneous zero-energy (low-energy) many-body states $|n(\phi)\rangle$, which are constructed as BdG Fock states

$$|n(\phi)\rangle \equiv \prod_{k=0}^{M-1} b_k^\dagger(\phi)^{n_k} |0(\phi)\rangle_b^p \quad (10.53)$$

through Bogoliubov diagonalisation of $H_{\text{Kit}}(\phi)$, as outlined in Secs. 5 and 5.3. Here, the $b_k^\dagger(\phi)$ are the (ϕ -dependent) creation operators of the M complex zero-energy (low-energy) Bogoliubov quasiparticle modes associated with the M pairs of MZMs in the system, $|0(\phi)\rangle_b^p$ denotes the product form of the (ϕ -dependent) Bogoliubov vacuum discussed in Sec. 5.3, and n_k is the occupation numbers of the k -th zero-energy (low-energy) Bogoliubov mode in $|n(\phi)\rangle$. In the following we will use a binary notation

$$n = \sum_{k=0}^{M-1} n_k 2^k \quad (10.54)$$

to enumerate our zero-energy (low-energy) states, i.e. $|0(\phi)\rangle = |0, \dots, 0\rangle$, $|1(\phi)\rangle = |1, 0, \dots, 0\rangle$ and so on. The non-Abelian statistics of the boundary MZMs is encoded in the non-Abelian Wilczek–Zee (WZ) phase matrix [197, 254]

$$\mathcal{U}_{\text{WZ}}(C) = \mathcal{P} \exp \left[- \oint_C \mathcal{A}_{d_0}(\phi) \right], \quad (10.55)$$

described in Sec. 4.3. Here, the non-Abelian Berry connection $A(\phi)$ has elements

$$\mathcal{A}_{d_0, mn}(\phi) = \langle m(\phi) | \partial_\phi | n(\phi) \rangle d\phi, \quad (10.56)$$

and the closed paths are continuous maps

$$C : [0, 1] \rightarrow \mathbb{S}_\phi^1, \quad t \mapsto \phi = C(t) \quad (10.57)$$

satisfying $C(0) = C(1)$. Numerically, we determine the WZ phase using the discretised version

$$\mathcal{U}_{\text{WZ}}(C) \approx \prod_{j=0}^{I-1} a_{d_0}(\phi_{I-j}) \quad (10.58)$$

of Eq. (10.55) that was introduced in Sec. 4.6. Here, the ϕ_j are taken from a discretised version

$$C_I = \{\phi_0, \phi_1, \dots, \phi_{I-1}, \phi_I = \phi_0\} \quad (10.59)$$

of the original continuous loop C , and the elements of the of the $d_0 \times d_0$ matrices $a_{d_0}(\phi_j)$ are BdG overlaps

$$a_{d_0, mn}(\phi_j) = \langle m(\phi_j) | n(\phi_{j-1}) \rangle, \quad (10.60)$$

which we compute by means of the Bertsch–Robledo formula Eq. (5.217) in Sec. 5.4. The parameter I controls the resolution of the discretised approximation.

Recall from Sec. 4.3 that the WZ phase matrix is generally gauge covariant, i.e. it transforms as

$$\mathcal{U}_{\text{WZ}}(C) \mapsto U^\dagger(C) \mathcal{U}_{\text{WZ}}(C) U(C) \quad (10.61)$$

under $U(C) \in \text{U}(d_0)$ gauge transformations

$$|n(\phi)\rangle \mapsto \sum_m |m(\phi)\rangle U_{mn}(\phi) \quad (10.62)$$

of the basis states $|n(\phi)\rangle$ of $\mathcal{H}_0(\phi)$. It is therefore important to address the “spectral structure” of the subspace $\mathcal{H}_0(\phi)$. Generally, two cases can be distinguished. In case (i) the energies of the eigenstates $|n(\phi)\rangle$ spanning $\mathcal{H}_0(\phi)$ are degenerate *for all* parameters $\phi \in \mathbb{S}_\phi^1$. Accordingly, $\mathcal{H}_0$ has the full $\text{U}(d_0)$ gauge symmetry and the only gauge-invariant observables are related quantities like the Wilson loop $\text{tr}(\mathcal{U}_{\text{WZ}}(C))$ or the determinant $\det(\mathcal{U}_{\text{WZ}}(C))$. In what follows, we are concerned with case (ii), where the energies of the eigenstates $|n(\phi)\rangle$ spanning $\mathcal{H}_0(\phi)$ are *non-degenerate almost everywhere* along $\mathbb{S}_\phi^1$ with potential degeneracies occurring only at isolated points. Hence, the physical gauge-freedom of $\mathcal{H}_0(\phi)$ is greatly reduced: the eigenstates are only fixed up to individual $\text{U}(1)$ phase factors, and the gauge group reduces to the maximal torus subgroup

$$\text{T} = \text{U}(1) \times \cdots \times \text{U}(1) < \text{U}(d_0). \quad (10.63)$$

In the minimal example with $d_0 = 2$, an element $T(C) \in \text{T}$ takes the form

$$T(C) = \begin{pmatrix} e^{i\alpha(C)} & 0 \\ 0 & e^{i\beta(C)} \end{pmatrix}, \quad (10.64)$$

and we obtain the transformation

$$T^\dagger(C) \mathcal{U}_{\text{WZ}}(C) T(C) = \begin{pmatrix} \mathcal{U}_{11}(C) & \mathcal{U}_{12}(C) e^{-i\gamma(C)} \\ \mathcal{U}_{21}(C) e^{i\gamma(C)} & \mathcal{U}_{22}(C) \end{pmatrix}, \quad (10.65)$$

where $\gamma(C) = \alpha(C) - \beta(C)$. The reduced gauge-freedom means that more properties of the WZ phase matrix can be physically distinguished. For instance, Eq. (10.65) shows that $\mathcal{U}_{\text{WZ}}(C) \simeq \mathbb{1}_2$, $\mathcal{U}_{\text{WZ}}(C) \simeq \sigma_z$ and $\mathcal{U}_{\text{WZ}}(C) \simeq \sigma_x$ are not gauge-equivalent. At the same time, $\mathcal{U}_{\text{WZ}}(C) \simeq \sigma_x$ is related to $\mathcal{U}_{\text{WZ}}(C) \simeq \sigma_y$ by a torus gauge-transformation.

Finally, note that the numerical implementation adopts a random gauge, so that all numerically obtained WZ phase matrices come out in a random gauge as well. For clarity, we clean up the numerical results and present gauge-fixed representatives in the following. Most notably, we are going to subsume all WZ matrices $\mathcal{U}$ with $|\mathcal{U}_{11}| = |\mathcal{U}_{22}| = 0$ and $|\mathcal{U}_{12}| = |\mathcal{U}_{21}| = 1$ under $\mathcal{U} \equiv \sigma_x$, since all (unitary) superpositions of σ_x and σ_y are gauge-equivalent under torus gauge-transformations of the form given in Eq. (10.65).

10.3 Exchangeless Braiding in One Kitaev Chain

In the next sections we apply the comprehensive many-body framework outlined in Sec. 10.2 to demonstrate non-Abelian exchangeless braiding in networks of finite-size Kitaev chains. Specifically, we aim to showcase the persistence of simple topological quantum gates in scenarios where finite interactions between boundary MZMs – arising from short chain lengths or significant inter-subchain couplings – lift the degeneracy of the many-body ground states.

As a starting point, we consider the Kitaev model on a finite chain of length L as described by H_{Kit} from Eq. (10.1). In the topologically non-trivial phase, the system hosts a pair of MZMs, Γ^A and Γ^B , localised at the left and right ends of the chain, respectively. As detailed in Sec. 6.4, the two MZM operators can be combined into a single complex zero-energy Bogoliubov operator

$$b_0 = \frac{1}{2} (\Gamma^A + i\Gamma^B), \quad (10.66)$$

which allows us to construct two orthogonal many-body ground states

$$|0\rangle_b \quad \text{and} \quad |1\rangle_b = b_0^\dagger |0\rangle_b, \quad (10.67)$$

where $|0\rangle_b \equiv |\bar{0}\rangle_b^P$ denotes the truncated product form of the quasiparticle vacuum defined in Eq. (5.183). Both $|0\rangle_b$ and $|1\rangle_b$ are eigenstates of the Bogoliubov quasiparticle number operator

$$N_b = \sum_m b_m^\dagger b_m, \quad (10.68)$$

satisfying

$$N_b |0\rangle_b = 0 \quad \text{and} \quad N_b |1\rangle_b = |1\rangle_b, \quad (10.69)$$

respectively. By extension, they are also eigenstates of the Bogoliubov quasiparticle parity operator

$$P_b = (-1)^{N_b}, \quad (10.70)$$

with which

$$P_b |0\rangle_b = (-1)^0 |0\rangle_b = +|0\rangle_b \quad \text{and} \quad P_b |1\rangle_b = (-1)^1 |1\rangle_b = -|1\rangle_b, \quad (10.71)$$

showing that $|0\rangle_b$ and $|1\rangle_b$ have opposite Bogoliubov quasi-fermion parities.⁴ The span

$$\mathcal{H}_0 = \text{span}\{|0\rangle_b, |1\rangle_b\} \quad (10.72)$$

of $|0\rangle_b$ and $|1\rangle_b$ defines the (almost) degenerate $d_0 = 2$ dimensional zero-energy (low-energy) subspace $\mathcal{H}_0$ that we are interested in going forward.

If the interaction Eq. (10.47) between the boundary MZMs Γ^A and Γ^B is sufficiently weak, they are projectively equivalent to Ising anyons [110, 111]. As a result, an adiabatic exchange of MZMs *in real space* induces a unitary transformation U_{AB} on the subspace $\mathcal{H}_0$ of degenerate many-body ground states that matches the Ising anyon braiding matrix [255]

$$R_{\text{Ising}} = e^{-i\pi/8} \begin{pmatrix} 1 & 0 \\ 0 & e^{i\pi/2} \end{pmatrix}, \quad (10.73)$$

up to a global phase factor, i.e.

$$U_{AB} = e^{i\alpha} R_{\text{Ising}}. \quad (10.74)$$

This was reviewed in Sec. 6.4. Here, we consider *exchangeless* braiding [206, 256]. With a complex superconducting (SC) gap parameter $\Delta = \Delta_0 e^{i\phi}$, the Kitaev-chain Hamiltonian in Eq. (10.1) defines a continuous family of Hamiltonians $H_{\text{Kit}}(\phi)$ parameterised by the phase $\phi \in \mathbb{S}_\phi^1$ of Δ . For given ϕ , the instantaneous many-body ground states $|0(\phi)\rangle_b$ and $|1(\phi)\rangle_b$ span the ϕ -dependent subspace $\mathcal{H}_0(\phi)$ from Eq. (10.72). The geometric evolution of $\mathcal{H}_0(\phi)$ along closed curves $C \subset \mathbb{S}_\phi^1$ is described by the WZ phase matrix $\mathcal{U}_{\text{WZ}}(C)$, see Eq. (10.55) and Sec. 4.3 for details. Below, we explicitly determine $\mathcal{U}_{\text{WZ}}(C^N)$ for closed curves

$$C^N : [0, 1] \rightarrow \mathbb{S}_\phi^1, \quad t \mapsto 2\pi N t, \quad (10.75)$$

that cover the $\mathbb{S}_\phi^1$ parameter manifold an integer number of N times. As explained in Sec. 10.2, we use a discretisation of $\mathcal{U}_{\text{WZ}}(C^N)$ to calculate it numerically. Concretely, we choose a resolution $I \gg 1$ and discretise C^N as

$$C_I^N = \{C_0^N, C_1^N, \dots, C_{I-1}^N, C_I^N\} = \{0, 2\pi N/I, \dots, 2\pi N(I-1)/I, 2\pi N\}. \quad (10.76)$$

⁴Note that $|0\rangle_b$ and $|1\rangle_b$ are both eigenstates of the Bogoliubov b -fermion parity operator $P_b = (-1)^{N_b}$ and the elementary c -fermion parity operator $P_c = (-1)^{N_c}$. In either case, $|0\rangle_b$ and $|1\rangle_b$ have *opposite* fermion parities. However, while $|0\rangle_b$ and $|1\rangle_b$ are eigenstates of the Bogoliubov quasi-fermion number operator N_b , cf. Eq. (10.69), they fail to be eigenstates of the elementary c -fermion number operator N_c , as discussed in Sec. 5.3 and, for instance, evident from Eqs. (5.183) and (9.39). Concretely, the Bogoliubov vacuum is a superposition of all (accessible) states of either even or odd c -parity – depending on whether or not there exist filled single-particle states in the Bloch–Messiah decomposition Eq. (5.182). While it is clear that $|0\rangle_b$ ($|1\rangle_b$) has even (odd) b -parity, it is therefore not per se evident whether $|0\rangle_b$ and $|1\rangle_b$ have even or odd c -parity. The only thing that is guaranteed is that they have opposite c -parities.

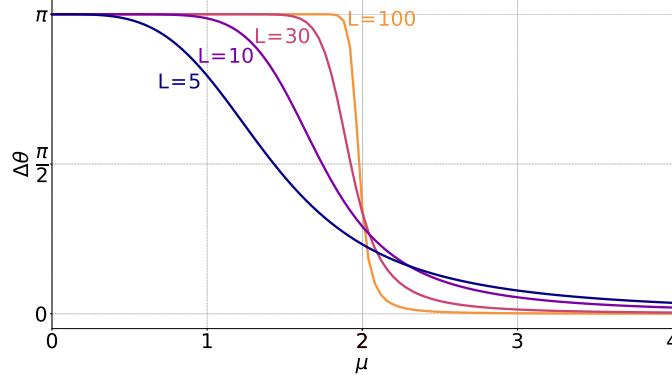


Figure 10.1: Difference $\Delta\theta$ between the phases acquired by $|0\rangle_b$ and $|1\rangle_b$ during a 2π rotation of the superconducting phase ϕ as a function of the chemical potential μ for different system sizes L . Chemical potential strengths of $\mu < 2$ ($\mu > 2$) correspond to the topologically non-trivial (trivial) phase. Parameters are $t = \Delta_0 = 1$. Adapted with minor modifications from Ref. [RQ3].

Based on this, we compute the 2×2 overlap matrices

$$a_2(\phi_j) = \begin{pmatrix} {}_b\langle 0(\phi_j)|0(\phi_{j-1})\rangle_b & {}_b\langle 0(\phi_j)|1(\phi_{j-1})\rangle_b \\ {}_b\langle 1(\phi_j)|0(\phi_{j-1})\rangle_b & {}_b\langle 1(\phi_j)|1(\phi_{j-1})\rangle_b \end{pmatrix} \quad (10.77)$$

between “adjacent” states $|n(\phi_j)\rangle_b$ and $|m(\phi_{j-1})\rangle_b$ along C_I^N using the Bertsch–Robledo formula given in Eq. (5.217). The path-ordered product Eq. (10.58) of matrices Eq. (10.77) then gives an approximation of $\mathcal{U}_{\text{WZ}}(C^N)$ the accuracy of which can be controlled using the resolution parameter I .

Let us first consider a Kitaev model with a “perfect” topological parameter configuration,

$$t = \Delta_0 = 1 \quad \text{and} \quad \mu = 0, \quad (10.78)$$

on a finite chain with $L = 20$ sites. As discussed towards the end of Sec. 10.1.1, the boundary MZMs of a perfect topological Kitaev chain are perfectly localised on the boundary sites and do not interact, regardless of the system size. Consequently, the energies of the two states $|0\rangle_b$ and $|1\rangle_b$ are exactly degenerate and the MZMs should exhibit flawless Ising anyon behaviour even for small chain lengths L . This makes “perfect” parameter configurations such as Eq. (10.78) particularly well-suited for benchmarking the numerical implementation. As a first test, we choose the simplest possible closed curve C^1 with winding number $N = 1$, and a resolution of $I = 10^3$. The latter is sufficient to reduce the element-wise numerical error of the WZ phase matrix to less than 1%. For the Kitaev chain defined by the parameter configuration in Eq. (10.78), the WZ phase matrix computed over the single-winding loop C^1 with resolution $I = 10^3$ is given by

$$\mathcal{U}_{\text{WZ}}(C^1) \approx e^{i\alpha} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} = \begin{pmatrix} e^{i\theta_0} & 0 \\ 0 & e^{i\theta_1} \end{pmatrix}. \quad (10.79)$$

Up to a global parameter-dependent phase factor, this corresponds precisely to the square of the Ising anyon braiding matrix from Eq. (10.73). Consequently, a 2π rotation of the superconducting phase in a topological Kitaev chain with anyonic boundary MZMs can be identified with a *double* exchange of these anyons in real space, realising a braiding process *without* physical exchange.⁵ This is what we call *exchangeless braiding*.

In the following, we will demonstrate that the form of $\mathcal{U}_{\text{WZ}}(C^1)$ in Eq. (10.79) is quite generic and robust against perturbations lifting the ground-state degeneracy. To this end, we first observe that the Ising anyon statistics of the MZMs is encoded in the difference

$$\Delta\theta(\mathcal{U}_{\text{WZ}}(C^1)) = (\theta_0 - \theta_1) \bmod 2\pi = \pi \quad (10.80)$$

⁵Accordingly, a closed path C^N , as given in Eq. (10.75), that implements N full rotations of the SC phase, generates a unitary evolution on the subspace $\mathcal{H}_0$ that can be understood as a $2N$ -fold exchangeless braiding process between Γ^A and Γ^B , if the system is in the topologically non-trivial phase and sufficiently large (depending on the model parameters t, Δ_0 and μ).

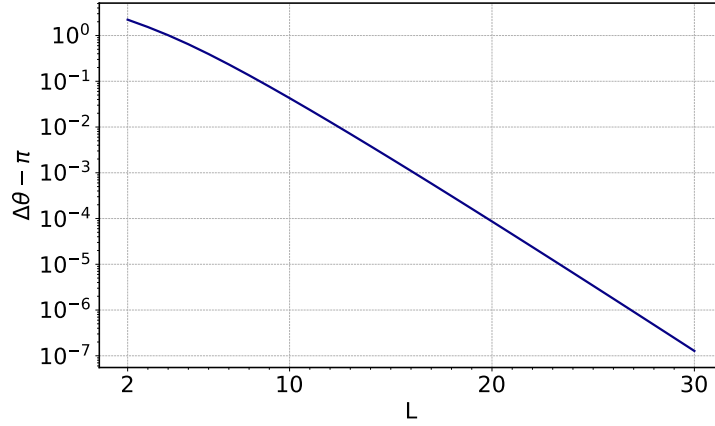


Figure 10.2: Logarithmic deviation of the phase difference $\Delta\theta$ from π , see Eq. (10.80), as a function of L for a single Kitaev chain. Parameters are $t = \Delta_0 = \mu = 1$. Adapted with minor modifications from Ref. [RQ3].

between the phases θ_0 and θ_1 acquired by the states $|0\rangle_b$ and $|1\rangle_b$ during the exchangeless braiding in Eq. (10.79).⁶ Thus, we may discard the physically insignificant global phase factor $e^{i\alpha}$ in Eq. (10.79) and focus our analysis on $\Delta\theta(U_{WZ}(C^1))$ for now.

We expect exchangeless braiding, as expressed by Eq. (10.80), to persist as long as the model supports topological Majorana boundary modes of sufficiently low energy. This requires that the Kitaev chain (i) remains in the non-trivial topological phase, and (ii) is long enough to suppress the interaction between the boundary MZMs.

To demonstrate (i), we fix the chain length L and vary the chemical potential μ . Since the nearest-neighbour hopping amplitude is set to $t = 1$, the topological class of the model is completely determined by the local potential μ : For $|\mu| < 2$ the system is topologically trivial; for $|\mu| > 2$ it is non-trivial. Accordingly, exchangeless braiding should be stable for any

$$-2 < \mu < 2. \quad (10.81)$$

Figure 10.1 shows the phase difference $\Delta\theta$ as function of the chemical potential μ for different values of L at $t = \Delta_0 = 1$. We find that $\Delta\theta$ gradually changes from $\Delta\theta = \pi$ to $\Delta\theta = 0$ as μ increases and passes the topological phase boundary at $\mu_c = 2$. The formation of a gradual transition regime is plausible since increasing μ simultaneously increases the MZM interaction strength K of the two-mode model Eq. (10.46) for any finite chain length L . This is corroborated by the observation that the width of the transition region decreases with increasing L , approaching a discontinuous step at $\mu_c = 2$ as $L \rightarrow \infty$. Notably, Fig. 10.1 also shows that even a relatively small system of $L = 30$ sites supports a substantial region in which $\Delta\theta = \pi$ is stable.

In order to address point (ii), we consider an “imperfect” topological parameter configuration

$$t = \Delta_0 = \mu = 1, \quad (10.82)$$

for which the decay parameters $x_{\pm}$ from Eq. (10.45), and hence the interaction length ℓ from Eq. (10.48) and the two-mode interaction strength K from Eq. (10.47), become non-trivial, so that the boundary MZMs experience a finite L -dependent interaction. In this setup, we proceed to vary the chain length L and track the deviation of $\Delta\theta$ from $\Delta\theta = \pi$. The result is shown Fig. 10.2. It suggests that $\Delta\theta - \pi$ can be made arbitrarily small by increasing L . In particular, even a chain of moderate length $L \approx 13$ is sufficient to get $|\Delta\theta - \pi| < 10^{-2}$ for the parameter configuration in Eq. (10.82). As $L \rightarrow \infty$, the deviation $|\Delta\theta - \pi|$ vanishes exponentially. This is consistent with the exponentially decreasing energy splitting between the states $|0\rangle_b$ and $|1\rangle_b$.

⁶It is worth mentioning that the opposite parities of $|0\rangle_b$ and $|1\rangle_b$ impose a superselection rule, so that ${}_b\langle 0|A|1\rangle_b = 0$ for all physical observables A . This superselection rule forbids coherent superpositions between $|0\rangle_b$ and $|1\rangle_b$ and makes it possible to understand the phases θ_0 and θ_1 from Eq. (10.79) as individual Berry phases even though the degenerate energies of $|0\rangle_b$ and $|1\rangle_b$ would generally require an analysis in terms of the $U(2)$ WZ phase.

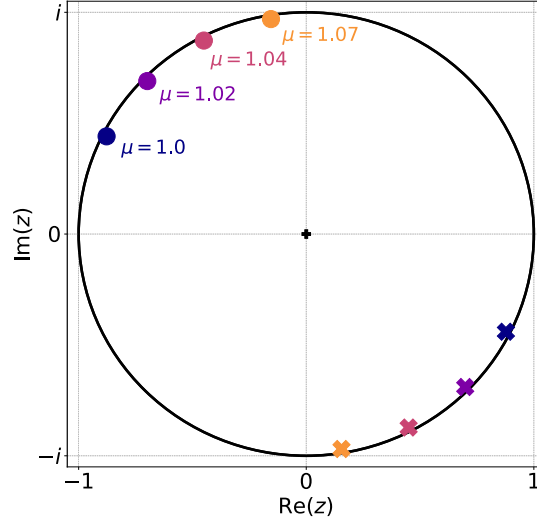


Figure 10.3: Individual geometric phases θ_0 (solid circles) and θ_1 (crosses) of ground state pairs $|0\rangle_b$ and $|1\rangle_b$ for different chemical potential strengths μ as indicated in the figure. Parameters are $t = \Delta_0 = 1$ and $L = 30$. Adapted with minor modifications from Ref. [RQ3].

It is worth mentioning that variations of L and μ within the stable region of Fig. 10.1 only affect the global phase α of U_{WZ} in Eq. (10.79), while preserving the anyonic phase difference of $\Delta\theta = \pi$. In particular, α scales linearly with L and, for large L , approaches the Sombbrero Berry phase

$$\alpha \approx \frac{L\pi}{2}, \quad (10.83)$$

that we determined for the Kitaev vacuum in Eq. (9.42) of Sec. 9.3. The μ -dependence of α close to $\mu = 1$ is shown in Fig. 10.3 for a Kitaev chain with $\Delta_0 = t = 1$ and $L = 30$. Again, the global phase α is very sensitive to small changes of μ and increases linearly with increasing μ .

The above considerations show that a single, spatially fixed Kitaev chain is enough to observe non-Abelian anyon physics. Despite this, it cannot be used directly for topological quantum computing (TQC). The reason for this is that qubits can only be realised in coherent two-dimensional Hilbert spaces, and the opposite fermion parities of the states $|0\rangle_b$ and $|1\rangle_b$ impose a superselection rule which precludes coherent superposition between them [102, 113]. As described in Sec. 6.3.2, the minimal setup for TQC with Ising anyons instead requires four anyonic MZMs. These give rise to a four-dimensional zero-energy (low-energy) space $\mathcal{H}_0$, whose two-dimensional even and odd parity sectors are then capable of accommodating a topological qubit.

10.4 Exchangeless Braiding in Two Kitaev Chains

One possibility to realise a minimal setup for TQC using Kitaev chains is to consider two separate chains that are weakly coupled at one boundary site. This is illustrated in Fig. 10.4 and captured by the Hamiltonian

$$H = H_1 + H_2 + H_W, \quad (10.84)$$

where H_1 and H_2 describe two Kitaev (sub-)chains of lengths L_1 and L_2 , see Eq. (10.1), and

$$H_W = -W \left(c_{L_1}^\dagger c_{L_1+1} + c_{L_1+1}^\dagger c_{L_1} \right) \quad (10.85)$$

implements a coupling of strength W between the last site $j = L_1$ of the first subchain and the first site $j = L_1 + 1$ of the second subchain. To aid the discussion, we define the total length $L = L_1 + L_2$ of the combined system and introduce an index $j = 1, \dots, L$ covering the sites of both subchains. Throughout

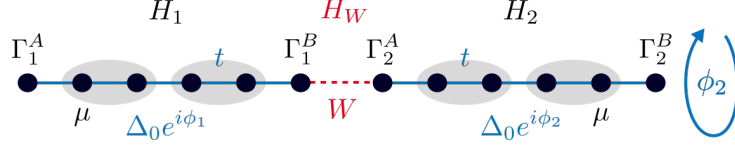


Figure 10.4: Sketch of two coupled Kitaev chains. Adapted with minor modifications from Ref. [RQ3].

this section, we use the same model parameters

$$t \equiv t_1 = t_2 = 1, \quad \Delta_0 \equiv \Delta_{0,1} = \Delta_{0,2}, \quad \mu \equiv \mu_1 = \mu_2 \quad (10.86)$$

for H_1 and H_2 , but allow different superconducting phases ϕ_1 and ϕ_2 . For configurations of the model parameters t, Δ_0 and μ that place H_1 and H_2 in the topologically non-trivial phase, the system features a total of four MZMs

$$\{\Gamma_1^A, \Gamma_1^B, \Gamma_2^A, \Gamma_2^B\}, \quad (10.87)$$

located on the four boundary sites

$$\{1, L_1, L_1 + 1, L\} \quad (10.88)$$

of the two subchains, as illustrated in Fig. 10.4. To ensure that these MZMs are well-defined, L_1 and L_2 must be large enough to suppress the interaction of MZMs within the subchains, and W must be weak enough to suppress the interaction of MZMs between them.

Analogous to the single Kitaev chain, Eq. (10.84) defines a continuous family of Hamiltonians $H(\phi_1, \phi_2)$ parameterised by the two independent SC phases $(\phi_1, \phi_2) \in \mathbb{S}_{\phi_1}^1 \times \mathbb{S}_{\phi_2}^1 \simeq \mathbb{T}_{\phi_1, \phi_2}^2$ of Δ_1 and Δ_2 . According to the scheme presented in Sec. 10.3, an exchangeless double braiding of the MZMs Γ_1^A and Γ_1^B of H_1 or Γ_2^A and Γ_2^B of H_2 can be induced by rotating the corresponding SC phase ϕ_1 or ϕ_2 by 2π . In preparation for a WZ-phase based analysis of such exchangeless braiding transformations, we extend the notion of the closed curves in $\mathbb{S}_{\phi}^1$ introduced in Eq. (10.75) to closed curves

$$C^{N_1, N_2} : [0, 1] \times [0, 1] \rightarrow \mathbb{S}_{\phi_1}^1 \times \mathbb{S}_{\phi_2}^1, \quad (t_1, t_2) \mapsto (2\pi N_1 t_1, 2\pi N_2 t_2) \quad (10.89)$$

in the parameter manifold $\mathbb{T}_{\phi_1, \phi_2}^2 = \mathbb{S}_{\phi_1}^1 \times \mathbb{S}_{\phi_2}^1$ of the two-chain system. The integers N_1 and N_2 indicate how often C^{N_1, N_2} winds around $\mathbb{S}_{\phi_1}^1 \subset \mathbb{T}_{\phi_1, \phi_2}^2$ and $\mathbb{S}_{\phi_2}^1 \subset \mathbb{T}_{\phi_1, \phi_2}^2$, respectively. In this notation, the WZ phase along the closed curve $C^{1,0}$ describes the exchangeless double braiding of the MZMs Γ_1^A and Γ_1^B of the first subchain, whereas the WZ phase along the closed curve $C^{0,1}$ captures the exchangeless double braiding of the MZMs Γ_2^A and Γ_2^B of the second subchain. As before, the numerical computation of the WZ phase is based on a resolution- I discretisation $C_I^{N_1, N_2}$ of C^{N_1, N_2} , which in this case is defined as the component-wise pairing (or *zipping*)

$$C_I^{N_1 N_2} \equiv C_I^{N_1} \times C_I^{N_2} := \{(C_0^{N_1}, C_0^{N_2}), (C_1^{N_1}, C_1^{N_2}), \dots, (C_{I-1}^{N_1}, C_{I-1}^{N_2}), (C_I^{N_1}, C_I^{N_2})\}, \quad (10.90)$$

where $C_I^{N_1}$ and $C_I^{N_2}$ take the form defined in Eq. (10.76). Importantly, the presence of four MZMs makes the effect of an exchangeless braiding process on $\mathcal{H}_0$ more nuanced. This can be understood as follows. While a system with precisely two MZMs Γ^A and Γ^B admits a single unique⁷ complex fermionic zero mode $b_0 = \frac{1}{2}(\Gamma^A + i\Gamma^B)$, there is an ambiguity in the definition of the two complex fermionic zero modes $b_{0,1}$ and $b_{0,2}$ supported by a system with four MZMs $\Gamma^A, \Gamma^B, \Gamma^C, \Gamma^D$, cf. Sec. 5.5. For instance, both

$$b_{0,1} = \frac{1}{2}(\Gamma^A + i\Gamma^B), \quad b_{0,2} = \frac{1}{2}(\Gamma^C + i\Gamma^D) \quad (10.91)$$

and

$$b'_{0,1} = \frac{1}{2}(\Gamma^A + i\Gamma^D), \quad b'_{0,2} = \frac{1}{2}(\Gamma^C + i\Gamma^B) \quad (10.92)$$

⁷Up to a complex phase factor, see Sec. 5.5.

represent valid choices. Either pair of complex fermion modes spans the same zero-energy (low-energy) subspace

$$\begin{aligned}\mathcal{H}_0 &= \text{span}\left\{|0,0\rangle_b, |1,0\rangle_b, |0,1\rangle_b, |1,1\rangle_b\right\} \\ &= \text{span}\left\{|0,0\rangle'_b, |1,0\rangle'_b, |0,1\rangle'_b, |1,1\rangle'_b\right\},\end{aligned}\quad (10.93)$$

where

$$\left\{|0,0\rangle_b, |1,0\rangle_b, |0,1\rangle_b, |1,1\rangle_b\right\} \equiv \left\{|0\rangle_b, b_{0,1}^\dagger |0\rangle_b, b_{0,2}^\dagger |0\rangle_b, b_{0,1}^\dagger b_{0,2}^\dagger |0\rangle_b\right\} \quad (10.94)$$

and

$$\left\{|0,0\rangle'_b, |1,0\rangle'_b, |0,1\rangle'_b, |1,1\rangle'_b\right\} \equiv \left\{|0\rangle_b, b_{0,1}'^\dagger |0\rangle_b, b_{0,2}'^\dagger |0\rangle_b, b_{0,1}'^\dagger b_{0,2}'^\dagger |0\rangle_b\right\}, \quad (10.95)$$

respectively. Accordingly, Eqs. (10.91) and (10.92) correspond to different basis choices of $\mathcal{H}_0$. This is important because the unitary transformation U_{XY} on $\mathcal{H}_0$ that is induced by exchanging any two

$$\Gamma^X \neq \Gamma^Y \in \{\Gamma^A, \Gamma^B, \Gamma^C, \Gamma^D\} \quad (10.96)$$

depends strongly on the way in which the four topological MZM operators $\Gamma^A, \Gamma^B, \Gamma^C, \Gamma^D$ combine into the two complex fermionic zero-mode operators $b_{0,1}$ and $b_{0,2}$ used in the construction of $\mathcal{H}_0$. In particular, the formal exchange of MZMs Γ^X and Γ^Y results in a qualitatively different unitary transformation depending on whether Γ^X and Γ^Y belong to the same complex fermion mode or not.

To illustrate this, consider a system with four MZMs $\Gamma^A, \Gamma^B, \Gamma^C, \Gamma^D$, which are paired into two complex fermionic zero-energy modes $b_{0,1}$ and $b_{0,2}$ as in Eq. (10.91). Following Eqs. (6.79) and (6.80) of Sec. 6.4, the unitary operator U_{AB} that exchanges Γ^A and Γ^B reads

$$U_{AB} = e^{\pi\Gamma^A\Gamma^B/4} = e^{i\pi/4} e^{-i\pi n_{0,1}/2}, \quad (10.97)$$

where $n_{0,1} = b_{0,1}^\dagger b_{0,1}$ is the quasiparticle number operator associated to $b_{0,1}$. We showed in Eqs. (6.75) and (6.76) of Sec. 6.4 that U_{AB} exchanges Γ^A and Γ^B like Ising anyons,

$$U_{AB}\Gamma^A U_{AB}^\dagger = -\Gamma^B \quad \text{and} \quad U_{AB}\Gamma^B U_{AB}^\dagger = \Gamma^A, \quad (10.98)$$

while transforming Γ^C and Γ^D trivially,

$$U_{AB}\Gamma^C U_{AB}^\dagger = \Gamma^C \quad \text{and} \quad U_{AB}\Gamma^D U_{AB}^\dagger = \Gamma^D. \quad (10.99)$$

Thus, the matrix representation $\mathcal{U}_{AB}$ of Eq. (10.97) with respect to the basis Eq. (10.94) is given by

$$\mathcal{U}_{AB} = \begin{pmatrix} e^{i\pi/4} & 0 & 0 & 0 \\ 0 & e^{-i\pi/4} & 0 & 0 \\ 0 & 0 & e^{i\pi/4} & 0 \\ 0 & 0 & 0 & e^{-i\pi/4} \end{pmatrix}, \quad (10.100)$$

cf. Eq. (6.94) of Sec. 6.4. Consequently, a double exchange of Γ^A and Γ^B is described by the diagonal $U(4)$ matrix

$$\mathcal{U}_{AB}^2 = e^{i\pi/2} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}. \quad (10.101)$$

If instead the four MZMs $\Gamma^A, \Gamma^B, \Gamma^C, \Gamma^D$ are paired into two complex fermionic zero-energy modes $b_{0,1}'$ and $b_{0,2}'$ as in Eq. (10.92), the unitary operator U_{AB} exchanging Γ^A and Γ^B becomes

$$U_{AB} = e^{\pi\Gamma^A\Gamma^B/4} = \frac{1}{\sqrt{2}} \left(1 + i(b_{0,1}'^\dagger + b_{0,1}') (b_{0,2}'^\dagger - b_{0,2}') \right), \quad (10.102)$$

see Eq. (6.100) of Sec. 6.4. The matrix representation $\mathcal{U}'_{AB}$ of Eq. (10.102) with respect to the basis Eq. (10.95) then reads

$$\mathcal{U}'_{AB} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & i & 0 & 0 \\ i & 1 & 0 & 0 \\ 0 & 0 & 1 & -i \\ 0 & 0 & -i & 1 \end{pmatrix}. \quad (10.103)$$

A double exchange of Γ^A and Γ^B in the alternative basis of $\mathcal{H}_0$ therefore results in the non-diagonal U(4) matrix

$$\mathcal{U}_{AB}^{\prime 2} = \begin{pmatrix} 0 & i & 0 & 0 \\ i & 0 & 0 & 0 \\ 0 & 0 & 0 & -i \\ 0 & 0 & -i & 0 \end{pmatrix}, \quad (10.104)$$

cf. Eqs. (6.101) and (6.103) of Sec. 6.4. In a system with perfectly degenerate MZMs, this basis dependence of the braiding outcome $\mathcal{U}_{XY}$ is compensated for by the covariance of $\mathcal{U}_{XY}$ under U(4) basis (gauge) transformations within $\mathcal{H}_0$: the braiding matrices of different basis choices transform into one another under the respective basis transformation. As was pointed out earlier, this changes when the degeneracy of the MZMs, and hence that of $\mathcal{H}_0$ as a whole, is lifted.

For a chain of fixed length $L < \infty$ and fixed $\Delta_0 = t = 1$, the ratio between the weak-link strength W and the local potential μ controls the relative strength of inter- and intra-subchain interactions between the MZMs. In the following, we show that this relative strength in turn determines how the MZMs combine into the complex fermionic low-energy modes. Specifically, dominant *intra-subchain* interaction leads to pairings

$$b_{0,1} = \frac{1}{2}(\Gamma_1^A + i\Gamma_1^B), \quad b_{0,2} = \frac{1}{2}(\Gamma_2^A + i\Gamma_2^B), \quad (10.105)$$

while dominant *inter-subchain* interaction instead produces

$$b'_{0,1} = \frac{1}{2}(\Gamma_1^A + i\Gamma_2^B), \quad b'_{0,2} = \frac{1}{2}(\Gamma_1^B + i\Gamma_2^A). \quad (10.106)$$

Varying the ratio between W and μ therefore allows us to tune between Eqs. (10.105) and (10.106), so that we can control whether the exchangeless braiding between Γ_s^A and Γ_s^B induced by a 2π rotation of the SC phase in the s -th subchain ($s = 1, 2$) results in a braiding transformation of Z -type, Eq. (10.101), or a braiding transformation of X -type, Eq. (10.104).

We first discuss the trivial case of two disconnected chains described by $W = 0$. Here, the MZMs of either subchain can only combine into a complex fermion mode locally, yielding complex fermionic zero modes of the form Eq. (10.105). A 2π phase rotation in either subchain induces an exchangeless braiding between MZMs from the same complex fermion, giving a U(4) WZ phase matrix of the form Eq. (10.101) resembling a Z -gate on either fixed-parity subspace, cf. Eqs. (6.97) and (6.98) of Sec. 6.4.

For a finite but sufficiently weak link $W > 0$ between the two subchains, we expect essentially the same result. In a generic system with subchains of different lengths, $L_1 \neq L_2$, there are two interaction strengths K_1 and K_2 associated with the intra-subchain interactions between the MZMs of either subchain, see Eqs. (10.45) – (10.48). It appears reasonable to expect that Z -type braiding persists with increasing W , provided the stronger of the two intra-subchain couplings K_1 and K_2 remains greater than W . Accordingly, a deviation from Z -type braiding is expected to occur roughly when

$$W \simeq K_{\max} \equiv \max\{K_1, K_2\}, \quad (10.107)$$

where K_1 and K_2 are defined as in Eq. (10.47). Only once W exceeds the dominant intra-subchain coupling, i.e. when $W > K_{\max}$, does the inter-subchain pairing from Eq. (10.106) become energetically favourable over the intra-subchain pairing of Eq. (10.105), and the expected braiding type changes from Z to X .

To test this numerically, we consider systems with fixed model parameters

$$t = \Delta_0 = 1, \quad \mu = 0.8, \quad (10.108)$$

but varying system sizes L and weak-link strengths W . Specifically, we choose to work with even system sizes L and minimally asymmetric subchain lengths

$$L_1 = L/2 - 1, \quad L_2 = L/2 + 1. \quad (10.109)$$

The latter ensures that L_1 and L_2 have the same parity⁸ and helps to avoid unwanted degeneracies. The chemical potential strength of $\mu = 0.8$ is chosen to remain clearly separated from both the perfect topological parameter configuration ($\mu = 0$) and the topological phase boundary ($\mu = 2$).⁹ With this, we select the closed curve $C^{0,1}$ as defined in Eq. (10.89) and compute the associated WZ phase matrices $\mathcal{U}_{\text{WZ}}(C^{0,1})$ for a range of values

$$10 \leq L \leq 30, \quad 10^{-5} \leq W \leq 10^{-1}. \quad (10.110)$$

A comparison between the resulting WZ phase matrices and the Z - and X -type braiding matrices from Eqs. (10.101) and (10.104) then allows for a detailed analysis of the expected transition between Z - and X -type quantum gates.

Analogous to the numerical treatment of the single Kitaev chain, the individual WZ phase matrices are obtained via the path-ordered product Eq. (10.58) of the 4×4 overlap matrices

$$a_4(\phi_j) = \begin{pmatrix} {}_b\langle 0(\phi_j)|0(\phi_{j-1})\rangle_b & {}_b\langle 0(\phi_j)|1(\phi_{j-1})\rangle_b & {}_b\langle 0(\phi_j)|2(\phi_{j-1})\rangle_b & {}_b\langle 0(\phi_j)|3(\phi_{j-1})\rangle_b \\ {}_b\langle 1(\phi_j)|0(\phi_{j-1})\rangle_b & {}_b\langle 1(\phi_j)|1(\phi_{j-1})\rangle_b & {}_b\langle 1(\phi_j)|2(\phi_{j-1})\rangle_b & {}_b\langle 1(\phi_j)|3(\phi_{j-1})\rangle_b \\ {}_b\langle 2(\phi_j)|0(\phi_{j-1})\rangle_b & {}_b\langle 2(\phi_j)|1(\phi_{j-1})\rangle_b & {}_b\langle 2(\phi_j)|2(\phi_{j-1})\rangle_b & {}_b\langle 2(\phi_j)|3(\phi_{j-1})\rangle_b \\ {}_b\langle 3(\phi_j)|0(\phi_{j-1})\rangle_b & {}_b\langle 3(\phi_j)|1(\phi_{j-1})\rangle_b & {}_b\langle 3(\phi_j)|2(\phi_{j-1})\rangle_b & {}_b\langle 3(\phi_j)|3(\phi_{j-1})\rangle_b \end{pmatrix}, \quad (10.111)$$

between “adjacent” BdG Fock states $|n(\phi_j)\rangle_b$ and $|m(\phi_{j-1})\rangle_b$ along a fine ($I \gg 1$) discretisation

$$C_I^{0,1} = \{(0, 0), (0, 2\pi/I), \dots, (0, 2\pi(I-1)/I), (0, 2\pi)\}, \quad (10.112)$$

of $C^{0,1}$ as defined in Eq. (10.90). Note that we have used the shorthand notation $\phi_j \equiv (\phi_j^1, \phi_j^2) \in C_I^{0,1}$ and adopted the binary naming convention

$$\{|0\rangle_b, |1\rangle_b, |2\rangle_b, |3\rangle_b\} \equiv \{|0, 0\rangle_b, |1, 0\rangle_b, |0, 1\rangle_b, |1, 1\rangle_b\} \quad (10.113)$$

from Eq. (10.54) to improve readability in Eq. (10.111). The individual overlaps between BdG Fock states in Eq. (10.111) are again computed using the Bertsch–Robledo formula given in Eq. (5.217).

The comparison between the WZ phase matrices $\mathcal{U} \equiv \mathcal{U}_{\text{WZ}}(C^{0,1})$ and the Z - and X -type quantum gates from Eqs. (10.101) and (10.104) is based on the Frobenius norm $\|A\| = \sqrt{\text{tr}(A^\dagger A)}$ of a matrix A , which we use to define the distance

$$d_Z(\mathcal{U}) = \|\mathcal{U} - |Z|\| \quad (10.114)$$

of $\mathcal{U}$ from Z , see Eq. (10.101), and the distance

$$d_X(\mathcal{U}) = \|\mathcal{U} - |X|\| \quad (10.115)$$

of $\mathcal{U}$ from X , see Eq. (10.104). Here, $|A|$ is the matrix with elements $|A_{ij}|$. With this, the transition between the Z - and the X -gate regimes can be quantified by the interpolation parameter

$$d(\mathcal{U}) = \frac{d_Z(\mathcal{U})}{\sqrt{d_X^2(\mathcal{U}) + d_Z^2(\mathcal{U})}}, \quad (10.116)$$

which yields $d(\mathcal{U}) = 0$ if $d_Z(\mathcal{U}) = 0$, i.e. $\mathcal{U} \simeq Z$, and $d(\mathcal{U}) = 1$ if $d_X(\mathcal{U}) = 0$, i.e. $\mathcal{U} \simeq X$. Additionally, we have to verify that in the deep Z - and X -gate phases, the WZ phase matrix, *including the phase factors of the matrix elements*, approaches the form given in Eqs. (10.101) and (10.104), respectively.

⁸In this case, parity refers to the number parity of an integer, i.e. to its property of being even or odd.

⁹The analysis was repeated for various values of $0 < \mu < 1$, yielding the same qualitative results.

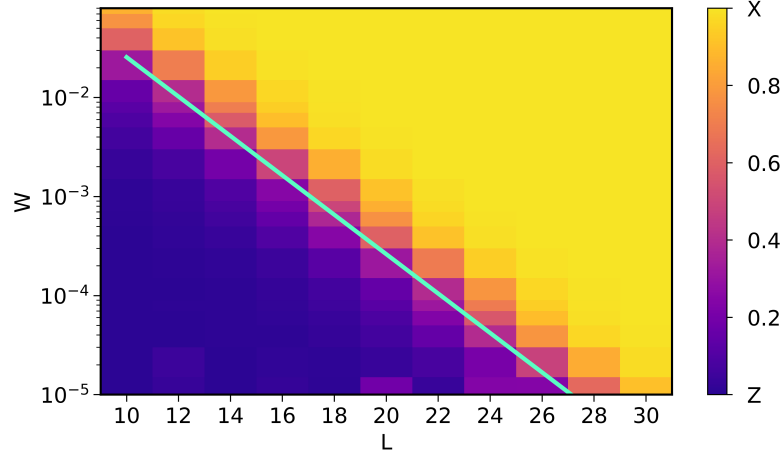


Figure 10.5: Phase diagram showing the braiding type of the WZ phase matrix as a function of the weak-link strength W and the total system size L for $\Delta_0 = t = 1$ and $\mu = 0.8$. The resolution I of the discretisation is set to $I = 1000$, which has proven sufficient to obtain converged results. The colour code shows $d = d(\mathcal{U}_{\text{WZ}}(C^{0,1}))$, see Eqs. (10.114) – (10.116) and surrounding discussion. In particular, we have $d = 0$ (purple) for a Z -type WZ phase, Eq. (10.120), and $d = 1$ (yellow) for an X -type WZ phase, Eq. (10.126). The cyan line corresponds to $W = K_{\text{max}} = 0.4^{L/2-1}$, cf. Eq. (10.119). Adapted with minor modifications from Ref. [RQ3].

Figure 10.5 shows the interpolation parameter $d = d(\mathcal{U}_{\text{WZ}}(C^{0,1}))$ from Eq. (10.116) evaluated for WZ phase matrices $\mathcal{U}_{\text{WZ}}(C^{0,1})$ of various two-Kitaev-chain systems with parameters as given in Eq. (10.108) and different combinations of weak link strengths W and system sizes L taken from the ranges specified in Eq. (10.110). The values of $d = d(\mathcal{U}_{\text{WZ}}(C^{0,1}))$ are represented as a colour gradient, ranging from purple ($d = 0$ and $\mathcal{U} \simeq Z$) to yellow ($d = 1$ and $\mathcal{U} \simeq X$). The “critical” weak link strength $W = K_{\text{max}}$, as obtained from Eqs. (10.107) and (10.47), is shown as a cyan line. By definition, $W = K_{\text{max}}$ decreases exponentially with increasing L . Specifically, the decay parameters $x_{\pm}$ from Eq. (10.45) evaluate to

$$x_{\pm}(1, \mu, 1) = \frac{-\mu \pm \mu}{4} \implies x_+ = 0 \quad \text{and} \quad x_- = -\frac{\mu}{2} \quad (10.117)$$

for topological parameter choices with $t = \Delta_0 = 1$ and $0 < \mu < 2$, such as Eq. (10.108). The corresponding characteristic length from Eq. (10.48) then becomes

$$\ell(1, \mu, 1) = \left| \ln \left(\frac{\mu}{2} \right) \right|^{-1}, \quad (10.118)$$

giving the interaction strength

$$K_{\text{max}}(1, \mu, 1) \propto \exp \left[-L_1 \left| \ln \left(\frac{\mu}{2} \right) \right| \right] = \left(\frac{\mu}{2} \right)^{L_1}, \quad (10.119)$$

where we plugged in $|\ln(\mu/2)| = -\ln(\mu/2)$ using $0 < \mu/2 < 1$, and where $L_1 = L/2 - 1$ is the shorter of the two subchain lengths. The latter appears as a result of the maximum in Eq. (10.107). Notably, $W = K_{\text{max}}$ also marks the approximate location of the transition between Z -type and X -type braiding, which is in line with the argument presented around Eq. (10.107). However, instead of a sharp transition between the Z - and X -gate phases at $W = K_{\text{max}}$, we observe a smooth crossover in its vicinity. In order to verify that the $d \simeq 0$ and $d \simeq 1$ regions in Fig. 10.5 actually correspond to Z - and X -type gates, we analyse the deep Z - and X -phases in greater detail.

We begin with the deep Z -gate phase for small L and weak W , see Fig. 10.5. There, we find WZ phase matrices of the form

$$\mathcal{U}_{\text{WZ}}(C^{0,1}) \approx e^{i\theta} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix} \equiv Z, \quad (10.120)$$

given in the basis $\{|0,0\rangle_b, |1,1\rangle_b, |0,1\rangle_b, |1,0\rangle_b\}$. Up to an irrelevant global parameter-dependent phase factor $e^{i\theta}$, this matches the expected result Eq. (10.101) for a system of disconnected subchains. Following Sec. 6.4, a restriction of the WZ phase matrix Eq. (10.120) to the even and odd parity sectors $\mathcal{H}_0|_{P_0}$ and $\mathcal{H}_0|_{P_1}$ of the (almost) degenerate subspace $\mathcal{H}_0$ yields

$$\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_0} = \mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_1} \approx e^{i\theta} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad (10.121)$$

where $\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_0}$ and $\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_1}$ are given in the bases $\{|0,0\rangle_b, |1,1\rangle_b\}$ and $\{|0,1\rangle_b, |1,0\rangle_b\}$, respectively. As for $W = 0$, see Eq. (10.80), we find

$$\Delta\theta(\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_0}) = \Delta\theta(\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_1}) \approx \pi. \quad (10.122)$$

for the phase differences of $\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_0}$ and $\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_1}$.

In the deep X -gate phase, i.e. for large L and strong W in Fig. 10.5, the numerical calculations yield WZ-phase matrices of the form

$$\mathcal{U}_{\text{WZ}}(C^{0,1}) \approx e^{i\theta} \begin{pmatrix} 0 & e^{-i\alpha} & 0 & 0 \\ e^{i\alpha} & 0 & 0 & 0 \\ 0 & 0 & 0 & e^{-i\beta} \\ 0 & 0 & e^{i\beta} & 0 \end{pmatrix}, \quad (10.123)$$

which is again given in the basis $\{|0,0\rangle_b, |1,1\rangle_b, |0,1\rangle_b, |1,0\rangle_b\}$. Up to a global parameter-dependent phase factor $e^{i\delta}$ and a suitable gauge-transformation

$$\mathcal{U}_{\text{WZ}}(C^{0,1}) \mapsto U \mathcal{U}_{\text{WZ}}(C^{0,1}) U^\dagger, \quad (10.124)$$

with a unitary matrix $U \in T < \text{U}(4)$ of the form

$$U = \text{diag}(e^{i\gamma_1}, e^{i\gamma_2}, e^{i\gamma_3}, e^{i\gamma_4}), \quad (10.125)$$

the WZ phase matrix $\mathcal{U}_{\text{WZ}}(C^{0,1})$ in Eq. (10.123) is equivalent to the expected result from Eq. (10.104). This can, for example, be seen using the global phase $\delta = \pi/2 - \theta$ and a gauge transformation U characterised by $\varphi_1 = -\varphi_2 = \alpha/2$ and $\varphi_3 = -\varphi_4 = (\beta - \pi)/2$, see App. A.12 for details. A similar calculation (for instance $\delta = -\theta$ combined with $\varphi_1 = -\varphi_2 = \alpha/2$ and $\varphi_3 = -\varphi_4 = \beta/2$) shows that Eq. (10.123) is also gauge-equivalent to the braiding transformation

$$\mathcal{U}_{\text{WZ}}(C^{0,1}) \approx e^{i\theta} \begin{pmatrix} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix} \equiv X. \quad (10.126)$$

Thus, the restrictions of the WZ phase matrix in Eq. (10.123) to the even and odd parity sectors $\mathcal{H}_0|_{P_0}$ and $\mathcal{H}_0|_{P_1}$ of $\mathcal{H}_0$ are both equivalent to

$$\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_0} = \mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_1} \simeq e^{i\theta} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad (10.127)$$

where $\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_0}$ and $\mathcal{U}_{\text{WZ}}(C^{0,1})|_{P_1}$ are given in the bases $\{|0,0\rangle_b, |1,1\rangle_b\}$ and $\{|0,1\rangle_b, |1,0\rangle_b\}$, respectively. For strong W and large L , the unitary transformation induced by exchangeless double braiding between the MZMs of either subchain can therefore be understood as an X -gate on the fixed-parity sectors of $\mathcal{H}_0$.

The above computation scheme enables an analysis of the transition between Z - and X -type WZ phase matrices as a function of the weak link strength W and the total system size L . Here, the total system size L effectively controls the intra-subchain interaction strengths K_1 and K_2 . Unsurprisingly, longer chains reduce intra-subchain interactions, favouring inter-subchain pairing and X -type braiding, while shorter chains enhance intra-subchain interactions, promoting intra-subchain pairing and Z -type braiding instead.

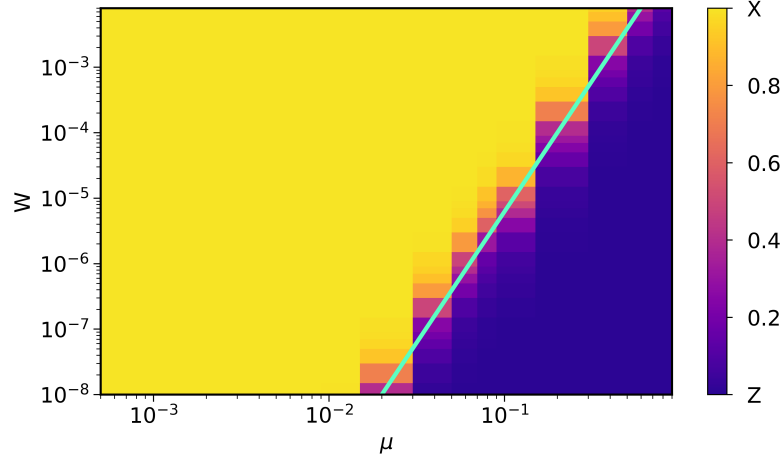


Figure 10.6: Phase diagram showing the braiding type of the WZ phase matrix as a function of the weak-link strength W and the chemical potential μ for $\Delta_0 = t = 1$ and $L = 10$. The resolution I of the discretisation is set to $I = 1000$, which has proven sufficient to obtain converged results. The colour code shows $d = d(\mathcal{U}_{\text{WZ}}(C^{0,1}))$ with $d = 0$ (purple) for a Z -type WZ phase, Eq. (10.120), and $d = 1$ (yellow) for an X -type WZ phase, Eq. (10.126). The cyan line corresponds to $W = K_{\text{max}} = (\mu/2)^4$, cf. Eq. (10.119) with $L_1 = 10/2 - 1 = 4$. Adapted with minor modifications from Ref. [RQ3].

However, Eqs. (10.45) – (10.48) show that the intra-subchain interaction strengths are not solely determined by the total system size L : variations of the other model parameters t , μ , and Δ_0 are expected to have a similar influence. In fact, the critical weak link strength $W = K_{\text{max}} = K_{\text{max}}(t, \mu, \Delta_0; L)$ defines a holonomic constraint on the parameter space spanned by (t, μ, Δ_0, L) , which gives rise to a critical hypersurface of codimension one. This is in line with Fig. 10.5, where we regard K_{max} as a function $K_{\text{max}} = K_{\text{max}}(W; L)$ of only two parameters W and L , giving the one-dimensional hypersurfaces shown as the cyan line. It seems natural to explore the impact of the other model parameters on the Z - to X -type braiding transition. Here, we focus on variations of the chemical potential strength μ . To this end, we repeat the previous analysis for fixed parameters

$$t = \Delta_0 = 1, \quad L = 10, \quad (10.128)$$

and a range of values

$$10^{-4} \leq \mu \leq 1, \quad 10^{-8} \leq W \leq 10^{-2}. \quad (10.129)$$

The resulting phase diagram is shown in Fig. 10.6. As expected from Eq. (10.119), weaker chemical potentials reduce intra-subchain interactions, facilitating inter-subchain pairing and X -type braiding, while stronger chemical potentials enhance intra-subchain interactions, favouring intra-subchain pairing and Z -type braiding. Although it is difficult to recognise in Fig. 10.6, the “critical” weak link strength $W = K_{\text{max}}$ varies as a power law $(\frac{2}{\mu})^{-L_1}$ with exponent $L_1 = L/2 - 1$, cf. Eq. (10.119).

Both, Figs. 10.5 and 10.6 show that the transition between the Z - and X -type braiding WZ phase matrices involves a crossover region, in which the braiding type of the WZ phase matrix is not defined. In either case, this region is quite accurately centered around the transition line $W = K_{\text{max}}$ (cyan lines) of the four-mode model discussed around Eq. (10.107). Moving away from the $W = K_{\text{max}}$ line – either towards weak W , small L , large μ , or towards strong W , large L , small μ – the system exits the finite crossover regime, and the WZ phase matrix rapidly approaches the Z - or X -type, respectively.

Importantly, Figs. 10.5 and 10.6 indicate that small systems of about $L = \mathcal{O}(10^1)$ sites are sufficient to realise Z or X quantum gates through exchangeless braiding. For instance, with $L = 12$ and $\mu = 0.8$, one can switch between Z and X by switching the weak-link strength between $W < 10^{-3}$ and $W > 10^{-1}$, see Fig. 10.5. For $L = 10$ and $\mu = 0.2$, a tiny change from $W \sim 10^{-6}$ to $W \sim 10^{-2}$ induces a transition from Z to X gates, see Fig. 10.6. We note that the location of the crossover parameter regime itself is reliably described by the four-mode model Eq. (10.107), whereas the width of the crossover is only captured by the WZ phase matrix of the full many-body theory.

11 – Conclusion and Outlook

This thesis aimed to advance our understanding of how geometric and topological properties inform the behaviour of quantum and quantum-classical systems. Through a sequence of case studies, we arrived at three key insights: (i) Coexisting topological structures can offer an explanation for the existence and form of spectral responses that quantum systems have to local impurities, (ii) helical boundary modes can reliably be harnessed to control the time-dependent state of magnetic impurities across mesoscopic length scales, and (iii) topological Majorana zero modes (MZMs) can be braided without physical exchange and retain most of their anyonic characteristics despite significant degeneracy-lifting perturbations.

Mathematical and Theoretical Preliminaries. To facilitate the overall discussion, the early chapters give a thorough review of the relevant concepts from mathematics and theoretical physics. With the goal of assembling a largely self-contained and physicist-friendly introduction to a technically demanding field, the mathematical part is reviewed from the ground up, while the physical material concentrates on specialised topics and methods. Notably, an overlap formula for Bogoliubov–de Gennes (BdG) vacua, originally proposed by Robledo in Ref. [92] and independently rediscovered during this work, is presented as part of this foundational segment in Sec. 5.4.

Spectral Impurity Responses in Systems with Coexisting Topological Structures. The first study [RQ1] explored how topological properties affect the spectral response of the spinful Haldane model to a small number of R magnetic impurities. By treating the magnetic impurities as classical spins of fixed lengths, their configuration space $\mathcal{S}_R$ is reduced to a $2R$ -dimensional parameter manifold $\mathcal{S}_R = \mathbb{S}_0^2 \times \cdots \times \mathbb{S}_{R-1}^2$, enabling a topological classification in terms of the R -th *spin*-Chern number $Ch_R^{(S)} \in \mathbb{Z}$. This spin-topological classification complements the conventional $\mathbf{k}$ -space topology of the Haldane model, which is classified by the first k -Chern number $C_1^{(k)} \in \mathbb{Z}$, characterising the bulk Bloch states over the two-dimensional Brillouin torus $\mathbb{T}_k^2$. The coexistence of the *intrinsic* $\mathbf{k}$ -space topology and the *extrinsic* $\mathcal{S}_R$ -space topology opens up novel perspectives and facilitates insights that are not readily attainable through either structure alone. Here, we use it to explain a robust spectral flow of single-particle energies that emerges when the exchange coupling J between the magnetic impurities and the Haldane host is varied from zero to infinity. This flow *bridges* the insulating gap of the Haldane host in the sense that for every chemical potential μ within the gap, there exists an exchange-coupling $J_{\text{crit}}(\mu)$ and a single-particle energy $\varepsilon_n(J)$ such that a single particle energy at this coupling is equal to the chemical potential, i.e. $\varepsilon_n(J_{\text{crit}}(\mu)) = \mu$. Importantly, the phenomenon occurs independently of the $\mathbf{k}$ -topological phase of the Haldane model. Its existence *can*, however, be linked to the $\mathcal{S}_R$ -topology of the hybrid model: at zero exchange-coupling strength the spin-Chern number vanishes trivially, giving $Ch_R^{(S)} = 0$, while for infinitiely strong coupling strength the system resembles a generalised magnetic monopole with spin-Chern number $Ch_R^{(S)} = 1$.

For a single magnetic impurity, this spin-topological phase transition implies the existence of some intermediate exchange-coupling strength $0 < J_{\text{crit}} < \infty$, at which the energy gap between the many-body ground state and a state with one more or less electron must close. This gap closure in turn requires a single-particle state $|\varepsilon_n(J)\rangle$ with $\varepsilon_n(J_{\text{crit}}) = \mu$. Since the same reasoning applies to every value of the chemical potential, the spin-topological phase transition provides a natural explanation for the observed spectral flow of single-particle energies bridging the gap. The spin-topology also gives a conceptual meaning to the critical exchange-interaction strength. Namely, it sets the radius of the magnetically charged two-sphere, which serves as the generalised (two-dimensional) magnetic monopole in the infinite coupling-strength limit. Since the total Hamiltonian is invariant under simultaneous $\text{SO}(3)$ rotations of the classical impurity spin and the quantum spin-degrees of freedom, the magnetic charge is distributed uniformly over the magnetic monopole sphere.

Our topological analysis of the spectral flow corroborates essential conclusions of earlier studies on spinless models with a single impurity [126, 131, 132, 135]. In particular, Slager *et al.* [131] highlight a diagnostic capacity of impurity-induced spectral responses in the context of local signatures of bulk topological order. Similar to their results [131] for a potential impurity in the $\mathbf{k}$ -topologically non-

trivial phase of the Bernevig–Hughes–Zhang (BHZ) model, we find that a magnetic impurity in the $\mathbf{k}$ -topologically non-trivial phase of the spinful Haldane model consistently induces a Zeeman pair of in-gap states whenever the exchange-coupling strength exceeds a finite threshold value $J_{\min}$. The energy splitting of the in-gap Zeeman pair decreases for increasing coupling strength and vanishes in the infinite coupling-strength limit, leaving two in-gap states with degenerate energies. In contrast, the $\mathbf{k}$ -topologically trivial phases of both the BHZ [131] and spinful Haldane models only support in-gap states within finite ranges of impurity strengths. Although the $\mathbf{k}$ -topology of the host systems does not affect the *existence* of the spectral flow, this difference shows that it still has a significant influence on its *qualitative* structure. In the trivial phase, the in-gap states quickly traverse the gap and vanish for $J \rightarrow \infty$, while in the non-trivial phase they enter the gap at $J_{\min}$ and remain inside the gap for $J \rightarrow \infty$. Consequently, the spectral flow for a $\mathbf{k}$ -topologically non-trivial host system only bridges the gap fully in the limit of $J \rightarrow \infty$. In this, our results for a magnetic impurity coupling to a single orbital in the unit cell of a spinful Haldane model are also in line with those obtained for a spinless impurity coupling to both orbitals in the unit cell of the spinless Haldane model [135]. In order to exploit this qualitative difference between *topological* and *trivial* spectral flows in practice, experimental methods are required to vary the potential strength. In Ref. [131], it was proposed that the impurity potential may be controlled experimentally by applying a tunable local gate voltage so that the presence or absence of in-gap states, and *thereby* the bulk topology, could be probed via scanning tunnelling spectroscopy. This applies to spin-resolved scanning tunneling microscope techniques as well [176]. Other experimental ways to realise and to control local impurities are discussed in Ref. [135].

The qualitative difference between the spectral flows of host systems with trivial and non-trivial $\mathbf{k}$ -space topology suggests an explanation rooted in bulk-boundary correspondence. Indeed, we were able to numerically verify an insightful argument for this. Based on the observation that for an infinitely strong coupling strength the impurity site is effectively removed from the host system, we interpret the $J \rightarrow \infty$ in-gap state as a super-discretised chiral edge mode living on the rudimentary boundary of the one-site impurity hole. This interpretation is supported by first considering a macroscopic hole, whose one-dimensional edge carries a genuine chiral boundary mode, and then gradually shrinking it to a single site. The resulting limiting process successively thins the edge-mode spectrum and links the dispersive chiral mode of the macroscopic hole to the super-discretised mode around the one-site impurity hole. In fact, even this super-discretised $J \rightarrow \infty$ in-gap impurity mode can be shown to carry a chiral current flowing around the impurity site. This analysis does not apply to the $\mathbf{k}$ -topologically trivial phase of the Haldane host system, explaining why the $J \rightarrow \infty$ in-gap modes persist only in the non-trivial phase.

The results for a spinful Haldane model with $R = 2$ impurity spins are similar to those obtained with $R = 1$ impurity spin. If the Haldane model is in the $\mathbf{k}$ -space topologically non-trivial phase, we observe at least one Zeeman pair of in-gap states, whose energies start out spin-split and become degenerate in the limit of infinite exchange coupling. For $J \rightarrow \infty$, these in-gap states can again be understood as super-discretised chiral edge modes. The occasional deviation from the naive expectation of *four* in-gap states for *two* holes can be explained geometrically: when the two impurity spins couple to *nearby* sites, the system may effectively form either one two-site hole (yielding two states), two one-site holes (yielding four states), or something in between, causing the number of pairs to *vary* with details of the electronic structure; at *larger separations*, the impurities behave as two distinct one-site holes, resulting in exactly two pairs. If the Haldane model is in the $\mathbf{k}$ -space topologically trivial phase, we find two in-gap modes fully bridging the bulk band gap within an intermediate range of coupling strengths. Again, this is explained by a spin-topological phase transition between $Ch_2^{(S)} = 0$ at $J = 0$, where the spin-configuration manifold S_2 is decoupled from the electronic system, and $Ch_2^{(S)} = 1$ for $J \rightarrow \infty$, where there are two magnetic monopoles with $Ch_1^{(S)} = 1$ each.

An important difference between the models with one and two impurity spins becomes apparent when the spin-topological transition is studied near the critical coupling strength J_{crit} . This results from an extensive numerical analysis of the critical exchange-coupling strengths in systems with a chemical potential in the middle of the bulk band gap. In the $\mathbf{k}$ -space topologically trivial case, we find that the transition takes place at a critical coupling of the order of the band width. In the non-trivial case, the critical coupling strength is typically larger because the in-gap modes correspond to super-discretised chiral impurity-boundary modes that converge to an energy within the band gap and therefore usually

only pass the mid-gap chemical potential at stronger couplings. In the exceptional case when the energies of the in-gap states converge exactly to μ for $J \rightarrow \infty$, there is no transition at all and we have $J_{\text{crit}} = \infty$. However, this scenario requires fine-tuning of the model parameters. Notably, the dependence of the typical magnitude of J_{crit} on the topological phase of the host system establishes a surprising correlation between the $\mathcal{S}_R$ -space and $\mathbf{k}$ -space topological phases. With two impurity spins, the spin-topological phase transition generally takes place in a finite range $J_{\text{crit},1} < J < J_{\text{crit},2}$ of coupling strengths. In this range the system is gapless on some subset of the spin configuration space $\mathcal{S}_2$ that depends on the coupling strength. Accordingly, the spin-topologically non-trivial phase is found for $J > J_{\text{crit},2}$, while the trivial phase is realised for $J < J_{\text{crit},1}$. In the transition range, the second spin-Chern number remains undefined. We also note that the gap closures at $J_{\text{crit},1}$ and $J_{\text{crit},2}$ take place for highly symmetric ferro/antiferromagnetic spin configurations.

Our findings point to a range of open problems and directions for future studies. The observed correlation between the $\mathbf{k}$ -space and $\mathcal{S}_R$ -space topological phases suggests that a similar connection might exist in other systems with non-trivial bulk topology, such as Chern insulators with k -Chern numbers greater than one, topological $\mathbb{Z}_2$ insulators, or topological superconductors. Moreover, it would be interesting to consider systems with a large number of $R \gg 1$ classical impurity spins and finally Kondo-lattice-type systems with $R \sim L$ as well. These also raise the question of how to compute high-order spin-Chern numbers $Ch_R^{(S)}$ with $R \gg 1$ in practice. Yet, even for small numbers of classical-spin impurities, there are interesting variations of our setup that are worth studying. This includes impurities with spin-anisotropic coupling, which reduces the symmetry of the gap-closure subset of the spin-configuration space. It might also be possible to realise and investigate larger spin-Chern numbers using impurity spins with short-range but non-local exchange couplings. Finally, methodological developments are needed to study bound states induced by *quantum*-spin impurities or impurities in correlated systems, including interacting topological insulators [257, 258].

Long-Range Helical Spin Control. While the first study [RQ1] concerned the *static* spectral response of a time-reversal symmetry (TRS) breaking Haldane model to magnetic impurities in the bulk, the second study [RQ2] addressed the *dynamics* of a magnetic impurity coupled to the helical boundary modes of a TRS invariant Kane–Mele model. The core motivation for this is simple: the protected edge modes of many topological quantum materials exhibit exotic and often desirable properties. While these properties have been studied extensively, their full potential for practical applications – especially at the microscopic level – continues to require detailed investigation.

Here, we have performed a comprehensive numerical analysis to demonstrate that the time-dependent state of a magnetic impurity can be controlled to a large extent by exploiting the helical character of the topological Kane–Mele edge modes. To this end, we have used a simple model consisting of a classical spin coupled to the zigzag edge of a Kane–Mele ribbon segment and studied various control protocols in detail. The key technical ingredient for the numerical treatment are dissipative Lindblad boundary conditions, which were set up to suppress reflections on all but the physically relevant edge of the ribbon segment. This allowed us to study the coupled microscopic real-time dynamics of the classical spin and the electronic system up to time scales of thousands of inverse hoppings without disturbing interference with reflected wave packets. In order to design a spin-switching protocol, we chained together several so-called *basic injection-pump* (BIP) processes. These consist of a dynamic spin injection at an injection site I of the physical zigzag edge and a subsequent pump phase during which the dynamics of the read-out spin at a site R of the physical zigzag edge is driven by the injected spin density. This process consists of four stages (A) – (D), all of which were designed to exploit the topological properties of the system:

(A) Local spin-up or spin-down excitations at an injection site I will induce a spin-polarised excitation that is predominantly carried by the edge states. The presence and helical character of these edge states is ensured by the bulk-boundary correspondence of the Kane–Mele model. Although the local magnetic field used to implement the injection also couples to the bulk states, the resulting bulk-state supported part of the injected spin-polarisation cloud is negligible and quickly dissipated into the bulk of the system. The remaining part of the spin-polarisation cloud is fully carried by the helical edge states and does not dissipate into the bulk. Importantly, we found that this edge-state supported spin-polarisation cloud saturates within a femtosecond time scale of a little more than 50 inverse hoppings. This makes the dynamic spin injection procedure a reproducible preparation step.

(B) By switching off the injection field, the spin excitation is released to propagate along the edge. Due to the helical nature and topological protection of the edge modes, this propagation is *unidirectional* and largely *lossless*. The group velocity of the propagating spin-polarisation cloud is naturally given by the Fermi velocity of the respective edge state. The fact that the edge-state propagation is mostly lossless ensures that the spin-polarisation cloud broadens, but does not lose volume. Consequently, the full volume of the injected spin-density cloud is available to drive the read-out spin dynamics even if the latter is located a mesoscopic distance away from the injection site.

(C) While TRS is intact in the pristine Kane–Mele ribbon segment, the classical spin at the read-out site R explicitly breaks TRS locally. This is important, because it facilitates a non-trivial interaction between the read-out spin and the propagating spin-polarisation cloud. Since TRS is largely restored away from the read-out site, the scattered spin-polarisation cloud is bound to split into two parts, namely one *onwards*-propagating part with the initial spin-polarisation, and one *backwards*-propagating part with the *opposite* spin polarisation. This means that the read-out spin exerts a finite spin torque on the polarisation cloud. The counter-torque, in turn, deflects the read-out spin towards the $\pm z$ -direction. The extent of this deflection depends on the initial alignment of the read-out spin. If the pump phase of the BIP process is sufficiently long, the process is reversible, and one may return to the initial spin state via the opposite spin-injection. We found that about 150 inverse hoppings, i.e. a few hundred femtoseconds, are needed to reach the fully relaxed, reversible final state.

(D) BIP processes can be concatenated to create a dynamical protocol. We have demonstrated that such a protocol can achieve a complete switching of the read-out spin orientation between the north and south poles within arbitrarily strict tolerances. Specifically, five to ten BIP processes are sufficient to switch a read-out spin within a tolerance of 1%. The full switching process can be inverted by reversing the polarisation of the spin in the dynamic injection part (A) of the BIP processes. Assuming a nearest-neighbor hopping energy of about 100 meV, the full switching process takes roughly a picosecond. However, this should be regarded as a lower limit. Since the system approaches a fully relaxed state after each BIP process, subsequent injection processes can be delayed. This implies that, in principle, the intermediate relaxed states during the switching process can be experimentally controlled even by techniques lacking pico- or sub-picosecond time resolution.

It is understood that our study describes an idealised model. A meaningful extension to real topological $\mathbb{Z}_2$ insulators must account for additional effects, including the impact of realistic multi-band electronic structures or electron-electron correlations. Another very interesting question is to which extent the conclusions remain valid for time-reversal-*symmetric* Kondo-type magnetic impurities, which can be modelled, for instance, by quantum spins [259, 260]. Systems featuring TRS-invariant Kondo impurities pose a challenging correlation problem, requiring treatment in the non-equilibrium regime and at long timescales. Even though the non-trivial topology is expected to largely protect the qualitative physics, it will also be interesting to examine the effects of surface imperfections and the specific placement and coupling-type of the magnetic impurities. More realistic descriptions of materials should also account for magnetic anisotropies, such as Dzyaloshinskii–Moriya interactions and Rashba spin-orbit coupling.

Exchangeless Braiding with Non-Degenerate Anyons. The third study [RQ3] presents a many-body analysis of exchangeless braiding that focusses on braiding-robustness against degeneracy-lifting hybridisations between anyons. In doing so, two major challenges in practical topological quantum computation (TQC) were addressed simultaneously. The first is the need for high-precision experimental control over physical anyon exchanges, which complicates exchange-based TQC schemes and is avoided here by adopting an *exchangeless* braiding protocol. The second concerns proximity-induced hybridisations between anyons. These inevitably arise in real-world experiments and are explicitly taken into account here by considering small systems featuring significant *degeneracy-lifting hybridisations*. Finally, the many-body treatment ensures that bulk contributions to the ground state are fully incorporated.

In order to explore hybridisation-impaired exchangeless braiding in a concrete and controllable environment, we considered a paradigmatic model system of weakly linked topological Kitaev chains. Specifically, we examined a single isolated Kitaev chain and a minimal network of two weakly-linked Kitaev chains. Since TQC with MZMs requires systems with at least four MZMs [101], we used the single-Kitaev-chain model (two MZMs) mainly as a proof of principle and regarded the two-Kitaev-chain network (four MZMs) as a minimal setup for TQC.

We started with the simple case of a single Kitaev chain of finite length. In the topological phase, the model features two boundary MZMs, which combine into a single complex Bogoliubov quasiparticle mode of (nearly) zero energy. This low-energy Bogoliubov quasiparticle mode gives rise to a two-dimensional subspace $\mathcal{H}_0$ of nearly degenerate many-body ground states. The MZMs of the topological Kitaev chain are known to be projectively equivalent to Ising anyons, and they are predicted to undergo an effective double exchange when the superconducting phase ϕ is advanced by 2π [14, 206]. Using the full many-body framework, we numerically confirmed this prediction by demonstrating that the $U(2)$ Wilczek–Zee (WZ) phase $U_{\text{WZ}}[0 \rightarrow 2\pi]$, which describes the geometric evolution of the two-dimensional low-energy many-body subspace $\mathcal{H}_0(\phi)$ throughout a full rotation of the superconducting phase ϕ , is projectively equivalent to the square of the Ising anyon exchange matrix. The main technical component for the numerical evaluation of the WZ phase is the Bertsch–Robledo formula [92], which enables an efficient computation of overlaps between BdG Fock states. Since the WZ phase can be reliably evaluated in networks of small chains with significant inter-chain coupling, we were able to confirm that the MZMs largely retain their anyonic properties despite the degeneracy-lifting hybridisations that occur under such conditions. Specifically, it is found that even short chains of about 30 sites support anyon physics across most of the topological phase. Surprisingly, it is precisely this deviation from the hybridisation-free ideal system that opens up access to a parameter regime with *tunable* braiding outcomes in the two-Kitaev chain network.

This hybridisation-supported regime emerges naturally upon extending the single-Kitaev-chain analysis to the minimal network of two weakly coupled Kitaev chains of finite length. The introduction of a second Kitaev chain has several immediate consequences. First, it doubles the number of MZMs, and with it the dimension of the low-energy subspace, from two to four. Additionally, it introduces two independent superconducting phases ϕ_1 and ϕ_2 , and three independent interaction strengths, namely the strength W of the weak link between the subchains and the two intra-subchain hybridisation strengths K_1 and K_2 that result from the finite lengths of the individual subchains. Upon studying the unitary WZ phase $U_{\text{WZ}}[0 \rightarrow 2\pi]$ describing the geometric evolution of the low-energy subspace $\mathcal{H}_0(\phi_1, \phi_2)$ under a full rotation of either superconducting phase ϕ_j , we find that the competition between W , K_1 , and K_2 controls a transition between two distinct parameter regimes: one in which $U_{\text{WZ}}[0 \rightarrow 2\pi]$ resembles a projective X gate, and one in which it resembles a projective Z gate. Specifically, we find a projective X gate when the weak-link strength dominates over the intra-subchain hybridisation strengths, and a projective Z gate when the intra-subchain hybridisation strengths dominate over the weak link.

Our study invites further exploration and development. We focussed on the non-Abelian WZ phase describing the geometry of the instantaneous low-energy many-body eigenstates. Although this is expected to constitute a tenable approximation in situations where the anyon braiding happens on timescales that are slow (adiabatic) compared to the superconducting gap, but fast (diabatic) compared to the hybridisation-induced energy-splitting between the low-energy many-body eigenstates, it will be interesting to repeat the many-body implementation of exchangeless braiding in networks of Kitaev chains using a time-dependent approach like the time-dependent BdG formalism [211]. Another very interesting possibility is to address the full quantum-classical real-time dynamics [172, 261] of more realistic hybrid systems. A particularly promising line of research concerns systems in which arrays of magnetic moments, modelled as classical spins, are locally exchange-coupled to a conventional s -wave superconductor to form so-called Yu-Shiba-Rusinov (YSR) chains. These YSR chains give rise to effectively one-dimensional dispersive YSR bands, which reside within the s -wave gap of the host system and may inherit a proximity-induced unconventional p -wave superconducting order. The spin-polarised p -wave superconducting YSR bands are then predicted to resemble Kitaev-chain physics. One expects that global rotations of the classical-spin chain configuration are equivalent to global rotations of the phase of the p -wave superconducting gap of the YSR bands [213]. In such quantum-classical models, the classical degrees of freedom used to control the exchangeless braiding process are themselves promoted to dynamic variables similar to those considered in [RQ2]. Moreover, the configuration space of these dynamic variables also provides a parameter manifold that enables an analysis of coexistent topological structures in the sense of [RQ1]. These setups therefore open up new fascinating problems at the interface between quantum-classical real-time dynamics, coexisting topological structures, and TQC with Ising MZMs.

A – Appendix

A.1 The Tenfold Way and Real Division Algebras

We are closely following [262] and materials [263]. Consider a Hilbert space $\mathcal{H}$. A normalised quantum state with representative $|\psi\rangle \in \mathcal{H}$ is an equivalence class

$$[\psi] = \left\{ e^{i\phi} |\psi\rangle, \phi \in \mathbb{R} \right\} \quad (\text{A.1})$$

of complex rays in $\mathcal{H}$. The space of these equivalence classes is called the projective Hilbert space $P(\mathcal{H})$. In this framework, a symmetry transformation is an automorphism

$$S : P(\mathcal{H}) \rightarrow P(\mathcal{H}), [\psi] \mapsto S[\psi] \quad (\text{A.2})$$

that preserves the ray product $[\phi] \cdot [\psi] := |\langle \phi | \psi \rangle|$ as

$$[\phi] \cdot [\psi] = [S\psi] \cdot [S\psi] \quad (\text{A.3})$$

between all states $[\phi], [\psi] \in P(\mathcal{H})$. Wigner's theorem states that every symmetry transformation $S : P(\mathcal{H}) \rightarrow P(\mathcal{H})$ comes either from a unitary operator $U_S : \mathcal{H} \rightarrow \mathcal{H}$ satisfying

$$\langle U_S \phi | U_S \psi \rangle = \langle \phi | U_S^\dagger U_S | \psi \rangle = \langle \phi | \psi \rangle \quad (\text{A.4})$$

or an antiunitary operator $A_S : \mathcal{H} \rightarrow \mathcal{H}$ satisfying

$$\langle A_S \phi | A_S \psi \rangle = \langle \phi | A_S^\dagger A_S | \psi \rangle^* = \langle \phi | \psi \rangle^* = \langle \psi | \phi \rangle. \quad (\text{A.5})$$

Suppose we have a symmetry transformation $S : P(\mathcal{H}) \rightarrow P(\mathcal{H})$ that squares to the identity as $S^2 = \mathbb{1}$. Such a symmetry transformation is called involutive. By Wigner's theorem, S comes either from a unitary transformation $U_S : \mathcal{H} \rightarrow \mathcal{H}$ with $U_S^2 = e^{i\phi} \mathbb{1}$ or an antiunitary transformation $A_S : \mathcal{H} \rightarrow \mathcal{H}$ with $A_S^2 = e^{i\phi} \mathbb{1}$. In the unitary case, we may multiply U_S by a phase,

$$U'_S := e^{-i\phi/2} U_S, \quad (\text{A.6})$$

to obtain a unitary operator U'_S that satisfies

$$U'^2_S := e^{-i\phi/2} U_S e^{-i\phi/2} U_S = e^{-i\phi} U_S^2 = e^{-i\phi} e^{i\phi} \mathbb{1} = \mathbb{1} \quad (\text{A.7})$$

and also corresponds to S . This is not possible in the antiunitary case, where

$$A'^2_S := e^{-i\phi/2} A_S e^{-i\phi/2} A_S = e^{-i\phi/2} e^{i\phi/2} A_S^2 = e^{i\phi} \mathbb{1} \quad (\text{A.8})$$

because of the antilinearity $A_S c = c^* A_S$ of A_S . However, combining $A_S e^{i\phi} = e^{-i\phi} A_S$ with

$$A_S e^{i\phi} = A_S A_S^2 = A_S^2 A_S = e^{i\phi} A_S \quad (\text{A.9})$$

yields the constraint

$$e^{i\phi} = e^{-i\phi} \implies e^{i\phi} = \pm 1. \quad (\text{A.10})$$

So if a symmetry $S : P(\mathcal{H}) \rightarrow P(\mathcal{H})$ with $S^2 = \mathbb{1}$ is implemented by a unitary operator U_S we can always choose U_S such that $U_S^2 = \mathbb{1}$, while an antiunitary implementation A_S of S always fulfills either $A_S^2 = +\mathbb{1}$ or $A_S^2 = -\mathbb{1}$. Now, one can show that an antiunitary implementation A_S with $A_S^2 = \mathbb{1}$ acts like a complex conjugation operator. That is, one can define a real Hilbert subspace

$$\mathcal{H}_{\mathbb{R}} := \{ |\psi\rangle \in \mathcal{H} \mid A_S |\psi\rangle = |\psi\rangle \} \subset \mathcal{H} \quad (\text{A.11})$$

containing only those vectors that are invariant under A_S . The original Hilbert space $\mathcal{H}$ is then simply the complexification

$$\mathcal{H} = \mathbb{C} \otimes_{\mathbb{R}} \mathcal{H}_{\mathbb{R}}. \quad (\text{A.12})$$

Given an antiunitary implementation A_S with $A_S^2 = -\mathbb{1}$ we can instead identify operators

$$I = i, \quad J = A_S, \quad K = IJ = iA_S \quad (\text{A.13})$$

that obey the quaternion relations

$$I^2 = J^2 = K^2 = IJK = -\mathbb{1}. \quad (\text{A.14})$$

Any given state $|\psi\rangle \in \mathcal{H}$ can therefore be multiplied by a quaternionic scalar $q \in \mathbb{H}$ as

$$q|\psi\rangle = (a + bI + cJ + dK)|\psi\rangle \quad (\text{A.15})$$

where $a, b, c, d \in \mathbb{R}$. Accordingly, we can understand $\mathcal{H}$ as a quaternionic Hilbert space $\mathcal{H}_{\mathbb{H}}$, i.e. a Hilbert space whose scalars come from the quaternions $\mathbb{H}$ instead of the complex numbers $\mathbb{C}$. The original Hilbert space $\mathcal{H}$ is then the underlying complex space

$$\mathcal{H} = \mathbb{C} \otimes_{\mathbb{C}} \mathcal{H}_{\mathbb{H}}. \quad (\text{A.16})$$

This demonstrates that quantum theory naturally accomodates instantiations of $\mathbb{R}$, $\mathbb{C}$ and $\mathbb{H}$ in the presence of involutive symmetry transformations. Mathematically, $\mathbb{R}$, $\mathbb{C}$ and $\mathbb{H}$ are special because they correspond to the only finite-dimensional associative division algebras over the real numbers (Frobenius theorem). A finite-dimensional associative algebra over the real numbers is a finite-dimensional real vector space V with an associative product $\cdot : V \times V \rightarrow V$ and a unit $\mathbb{1} \in V$. Such an algebra is called a division algebra if every non-zero element $v \in V$ has a multiplicative inverse $v^{-1} \in V$ with which $v^{-1} \cdot v = v \cdot v^{-1} = \mathbb{1}$.

Now consider a group G of symmetry transformations and a Hilbert space $\mathcal{H}$. A unitary representation ρ of G on $\mathcal{H}$ consists of unitary operators $\rho(g) : \mathcal{H} \rightarrow \mathcal{H}$ that satisfy

$$\rho(gh) = \rho(g)\rho(h) \quad \text{and} \quad \rho(e) = \mathbb{1}. \quad (\text{A.17})$$

We call ρ irreducible if the only invariant subspaces, i.e. the only subspaces $I \subset \mathcal{H}$ for which $\rho(g) : I \rightarrow I$ for all $g \in G$, are the trivial subspace $I = \{\mathbb{1}\}$ and the entire Hilbert space $I = \mathcal{H}$. Schur's lemma states that the only unitary operators $U : \mathcal{H} \rightarrow \mathcal{H}$ that commute as

$$U\rho(g) = \rho(g)U \quad (\text{A.18})$$

with all $\rho(g)$ are of the form

$$U = e^{i\phi} \mathbb{1}. \quad (\text{A.19})$$

Dyson showed that if ρ is an irreducible unitary representation of a symmetry group G on a Hilbert space $\mathcal{H}$, then precisely one of the following holds:

1. There exists an antiunitary operator A with $A^2 = -\mathbb{1}$ commuting with all $\rho(g)$. This endows $\mathcal{H}$ with a quaternionic structure and we call ρ a quaternionic representation.
2. There is no antiunitary operator A commuting with all $\rho(g)$ and we call ρ a complex representaiton.
3. There exists an antiunitary operator A with $A^2 = +\mathbb{1}$ commuting with all $\rho(g)$. This endows $\mathcal{H}$ with a real structure and we call ρ a real representation.

This classification scheme is known as Dyson's threefold way and it is clear from the previous considerations that it corresponds to the threefold way of real division algebras. The threefold way has profound implications for physics. For instance, consider the spin- n representations of $\text{SU}(2)$ that we use to describe the spin of (elementary) particles. These correspond to the $(2n + 1)$ -dimensional irreducible complex representation of $\text{SU}(2)$. For all spin- n representations of $\text{SU}(2)$ there exists an antiunitary involutive symmetry transformation A that commutes with all $\text{SU}(2)$ transformations. Now, when n is

T	C	S	Super Division Algebra
0	0	0	$\text{Cl}_0(\mathbb{C}) \simeq \mathbb{C}$
0	0	1	$\text{Cl}_1(\mathbb{C}) \simeq \mathbb{C} \oplus \mathbb{C}$
1	0	0	$\text{Cl}_{0,0}(\mathbb{R}) \simeq \mathbb{R}$
1	-1	1	$\text{Cl}_{0,1}(\mathbb{R}) \simeq \mathbb{C}$
0	-1	0	$\text{Cl}_{0,2}(\mathbb{R}) \simeq \mathbb{H}$
-1	-1	1	$\text{Cl}_{0,3}(\mathbb{R}) \simeq \mathbb{H} \oplus \mathbb{H}$
-1	0	0	$\text{Cl}_{0,4}(\mathbb{R}) \simeq M_2(\mathbb{H})$
-1	1	1	$\text{Cl}_{3,0}(\mathbb{R}) \simeq M_2(\mathbb{C})$
0	1	0	$\text{Cl}_{2,0}(\mathbb{R}) \simeq M_2(\mathbb{R})$
1	1	1	$\text{Cl}_{1,0}(\mathbb{R}) \simeq \mathbb{R} \oplus \mathbb{R}$

Table A.1: The symmetry classes of the tenfold way and the corresponding Morita equivalence classes of super division algebras. The symbol $\simeq$ means Morita equivalent.

an integer, A satisfies $A^2 = \mathbb{1}$ and the spin- n representation on $\mathbb{C}^{2n+1}$ is the complexification of a real representation on $\mathbb{R}^{2n+1}$. If n is a half-integer, A fulfils $A^2 = -\mathbb{1}$ and the spin- n representation on $\mathbb{C}^{2n+1}$ is the underlying complex representation of a quaternionic representation on $\mathbb{H}^{(2n+1)/2}$.

The tenfold way extends Dyson's threefold way based on a construction that is known as super Hilbert spaces. A super Hilbert space $\mathcal{H}$ is simply a Hilbert space can be written as a direct sum

$$\mathcal{H} = \mathcal{H}_0 \oplus \mathcal{H}_1 \quad (\text{A.20})$$

of an “even” Hilbert space $\mathcal{H}_0$ and an “odd” Hilbert space $\mathcal{H}_1$. There are many ways in which super Hilbert spaces may arise. For example, one can choose $\mathcal{H}_0$ to accomodate bosonic states and $\mathcal{H}_1$ to accomodate fermionic states. In condensed matter theory, we can write the Fock space $\mathcal{F}$ as a super Hilbert space

$$\mathcal{F} = \mathcal{F}_0 \oplus \mathcal{F}_1 \quad (\text{A.21})$$

consisting of the particle Hilbert space $\mathcal{F}_0$ and the hole Hilbert space $\mathcal{F}_1$. In this setting, we can have antiunitary operators $\mathcal{T} : \mathcal{F} \rightarrow \mathcal{F}$ that are even in the sense that they map

$$\mathcal{T} : \mathcal{F}_0 \rightarrow \mathcal{F}_0 \quad \text{and} \quad \mathcal{T} : \mathcal{F}_1 \rightarrow \mathcal{F}_1 \quad (\text{A.22})$$

and antiunitary operators $\mathcal{C} : \mathcal{F} \rightarrow \mathcal{F}$ that are odd in the sense that they map

$$\mathcal{C} : \mathcal{F}_0 \rightarrow \mathcal{F}_1 \quad \text{and} \quad \mathcal{C} : \mathcal{F}_1 \rightarrow \mathcal{F}_0. \quad (\text{A.23})$$

For physical reasons we call $\mathcal{T}$ the time-reversal symmetry (TRS) operator and $\mathcal{C}$ the particle-hole symmetry (PHS) operator. In a given condensed matter system we may then have $\mathcal{T}$ -symmetry with $\mathcal{T}^2 = \pm\mathbb{1}$ or no $\mathcal{T}$ -symmetry. Independently, we may have $\mathcal{C}$ -symmetry with $\mathcal{C}^2 = \pm\mathbb{1}$ or no $\mathcal{C}$ -symmetry as well. This gives rise to nine out of ten Altland-Zirnbauer symmetry classes [26]. The final class corresponds to systems that are neither TRS nor PHS symmetric, but still symmetric under the combined unitary transformation $\mathcal{S} = \mathcal{T} \cdot \mathcal{C}$ called chirality.

Just like Dyson's threefold way corresponds to the set of distinct associative real division algebras, the tenfold way corresponds to the set of distinct associative real super division algebras, of which there are precisely ten. Surprisingly, all of these are Morita¹ equivalent to complex or real Clifford algebras. This is rather exciting for a number of reasons. A physical one is that representations of real Clifford algebras are used to describe the spinful elementary particle fields so they are ubiquitous in the standard model and beyond. On the mathematical side, this allows for an intertwining of concepts: the Clifford algebras give Lie groups and symmetric spaces, establishing connections to Riemannian geometry and classifying spaces for vector bundles.

¹Essentially, two superalgebras A and B are called Morita equivalent, $A \simeq B$, if they have equivalent representations on super vector spaces.

A.2 The Real Schur Decomposition Theorem

To prove the real Schur decomposition theorem we first prove a lemma identifying the eigenspace of a pair of mutually conjugate complex eigenvalues of a real matrix as an invariant two-dimensional subspace. A subspace V of $\mathbb{R}^n$ is said to be invariant with respect to a matrix A , and we write $AV = V$, if for any $\mathbf{x} \in V$ we have that $A\mathbf{x} \in V$. Let A be a real $n \times n$ matrix with a complex eigenvalue $\lambda_1 = a + ib$ (with $b \neq 0$) and let $\mathbf{z} = \mathbf{x} + i\mathbf{y}$ be an eigenvector belonging to λ_1 . Then $\text{span}(\mathbf{x}, \mathbf{y}) =: V$ fulfills $\dim_{\mathbb{R}}(V) = 2$ and $AV = V$, i.e. V is a two-dimensional subspace of $\mathbb{R}^n$ that is invariant under A . This can be seen as follows. Since we assumed λ to be complex with non-zero imaginary part b the imaginary part $\mathbf{y}$ of its associated eigenvector $\mathbf{z}$ has to be non-zero as well. The fact that A is a real matrix tells us that $\bar{\lambda} = a - ib$ is also an eigenvalue of A and that $\bar{\mathbf{z}} = \mathbf{x} - i\mathbf{y}$ is an eigenvector belonging to $\bar{\lambda}$. Now we can use that $\mathbf{z}$ and $\bar{\mathbf{z}}$ are eigenvectors belonging to *distinct* eigenvalues λ and $\bar{\lambda}$ to show that there exists no complex scalar c such that $\mathbf{x} = c\mathbf{y}$ because if there was such a scalar c we could write $\mathbf{z} = (c + i)\mathbf{y}$ and $\bar{\mathbf{z}} = (c - i)\mathbf{y}$ which would make $\mathbf{z}$ and $\bar{\mathbf{z}}$ linearly dependent thus contradicting the guaranteed linear independence of eigenvectors belonging to distinct eigenvalues. Therefore $\mathbf{x}$ and $\mathbf{y}$ are linearly independent and span a two-dimensional subspace $V = \text{span}(\mathbf{x}, \mathbf{y})$ of $\mathbb{R}^n$. In order to show that V is invariant under A we note that $A\mathbf{z} = \lambda\mathbf{z}$ gives

$$A\mathbf{x} + iA\mathbf{y} = (a\mathbf{x} - b\mathbf{y}) + i(b\mathbf{x} + a\mathbf{y}) \quad (\text{A.24})$$

and hence

$$A\mathbf{x} = a\mathbf{x} - b\mathbf{y} \quad \text{and} \quad A\mathbf{y} = b\mathbf{x} + a\mathbf{y}. \quad (\text{A.25})$$

If $\mathbf{w} = k\mathbf{x} + l\mathbf{y}$ is any vector in V , then

$$A\mathbf{w} = kA\mathbf{x} + lA\mathbf{y} = kax - kby + lbx + lay = (ka + lb)\mathbf{x} + (la - kb)\mathbf{y} \quad (\text{A.26})$$

is clearly just another vector in V .

With this lemma at hand we can prove the real Schur decomposition as stated before by induction. Let $n = 2$. If the eigenvalues of A are real we can take $\mathbf{q}_1$ to be a unit eigenvector of the first real eigenvalue λ_1 and then $\mathbf{q}_2$ to be any unit vector that is orthogonal to $\mathbf{q}_1$. Then the matrix $Q = (\mathbf{q}_1, \mathbf{q}_2)$ is orthogonal by construction and the matrix $T := Q^T A Q$ is in upper triangular form with 1×1 eigenvalue blocks on its diagonal. This can be seen writing down the first column of T as

$$Q^T A \mathbf{q}_1 = \lambda_1 Q^T \mathbf{q}_1 = \lambda_1 \mathbf{e}_1, \quad (\text{A.27})$$

which immediately yields the upper triangular form as well as the first 1×1 eigenvalue block on its diagonal. The second diagonal entry of T necessarily has to be the second eigenvalue of A because we can Laplace expand the determinant of the characteristical polynomial with respect to the first column of T and get

$$\begin{aligned} 0 &\stackrel{!}{=} \det(T - \lambda \mathbb{1}_2) \\ &= \det \begin{pmatrix} \lambda_1 - \lambda & T_{12} \\ 0 & T_{22} - \lambda \end{pmatrix} \\ &= (\lambda_1 - \lambda)(T_{22} - \lambda). \end{aligned} \quad (\text{A.28})$$

Of course, Laplace expansion is rather overkill for $n = 2$ but we will use it as part of the induction so bear with me. If the eigenvalues of A are complex we may set $T = A$ and $Q = \mathbb{1}_2$ and be done. This last bit may seem like a trivial statement, and it somewhat is, but it still is an important point to make. The upper triangular block Schur form features 1×1 blocks on its diagonal for its real eigenvalues and 2×2 blocks for its pairs of mutually conjugate complex eigenvalues. In case of a 2×2 matrix A this means that we only have to do stuff if the two eigenvalues of A are real because only then do we have to massage A into the Schur form with 1×1 eigenvalue blocks on its diagonal and all that. If the eigenvalues of a 2×2 matrix A are a pair of mutually conjugate complex numbers the Schur decomposition is not going to do anything to A because A itself already is a 2×2 block on its diagonal, and therefore already has

the Schur upper triangular block form. So in either case we end up with an upper triangular block matrix of the Schur type for our $n = 2$ induction initial case. Now let A be a real $k \times k$ matrix with $k \geq 3$ and assume that every real $m \times m$ matrix with $2 \leq m \leq k$ has a Schur decomposition. Let λ_1 be an eigenvalue of A . If λ_1 is real, let $\mathbf{q}_1$ be a uni eigenvector of λ_1 and choose $(k-1)$ unit vectors $\mathbf{q}_2, \dots, \mathbf{q}_k$ spanning the orthogonal complement of $\mathbf{q}_1$ in $\mathbb{R}^k$ then $Q_1 = (\mathbf{q}_1, \dots, \mathbf{q}_k)$ is an orthogonal $k \times k$ matrix. As in the initial case it follows that the first column of $T_1 = Q_1^T A Q_1$ is equal to $\lambda_1 \mathbf{e}_1$ giving

$$T_1 = \begin{pmatrix} \lambda_1 & X \\ 0 & A_1 \end{pmatrix} \quad (\text{A.29})$$

with a real $(k-1) \times (k-1)$ matrix A_1 and some $(k-1)$ row vector X . By our induction hypothesis A_1 has a Schur decomposition $T_2 := O_2^T A_1 O_2$ with a $(k-1) \times (k-1)$ orthogonal matrix O_2 . We may define the $k \times k$ orthogonal matrix

$$Q_2 := \begin{pmatrix} 1 & 0 \\ 0 & O_2 \end{pmatrix} \quad (\text{A.30})$$

which together with Q_1 transforms A as

$$\begin{aligned} T &:= Q_2^T Q_1^T A Q_1 Q_2 \\ &= Q_2^T T_1 Q_2 \\ &= \begin{pmatrix} 1 & 0 \\ 0 & O_2^T \end{pmatrix} \begin{pmatrix} \lambda_1 & X \\ 0 & A_1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & O_2 \end{pmatrix} \\ &= \begin{pmatrix} \lambda_1 & X O_2 \\ 0 & O_2^T A_1 O_2 \end{pmatrix} \\ &= \begin{pmatrix} \lambda_1 & X O_2 \\ 0 & T_2 \end{pmatrix} \end{aligned} \quad (\text{A.31})$$

which has the desired upper triangular block form of the Schur decomposition. This proves that the $k \times k$ orthogonal transformation $Q := Q_1 Q_2$ brings A into Schur block form T , i.e. $T := Q^T A Q$. If λ_1 is complex, the argument proceeds analogously: we find an eigenvector $\mathbf{z} = \mathbf{x} + i\mathbf{y}$ and compute the span $V := \text{span}(\mathbf{x}, \mathbf{y})$. In the preceding lemma we convinced ourselves that V is a two-dimensional subspace of $\mathbb{R}^k$ and that V is invariant under A . Next we choose any orthonormal basis $\{\mathbf{q}_1, \mathbf{q}_2\}$ of V along with $(k-2)$ unit vectors spanning the orthogonal complement of V in $\mathbb{R}^k$. Then the matrix $Q_1 = (\mathbf{q}_1, \mathbf{q}_2, \dots, \mathbf{q}_k)$ is again an orthogonal matrix and the first two columns of $T_1 = Q_1^T A Q_1$ are

$$Q_1^T A \mathbf{q}_1 = B_{11} \mathbf{q}_1 + B_{21} \mathbf{q}_2 \quad \text{and} \quad Q_1^T A \mathbf{q}_2 = B_{12} \mathbf{q}_1 + B_{22} \mathbf{q}_2, \quad (\text{A.32})$$

due to the invariance $AV = V$ of V under A . We therefore find

$$T_1 = \begin{pmatrix} B & X \\ 0 & A_1 \end{pmatrix} \quad \text{where} \quad B = \begin{pmatrix} B_{11} & B_{12} \\ B_{21} & B_{22} \end{pmatrix} \quad (\text{A.33})$$

is the coefficient matrix from Eq. (A.32) and A_1 is now a $(k-2) \times (k-2)$ matrix.² Now we can apply our induction hypothesis on A_1 and proceed in exactly the same way as before to show that A may be Schur decomposed. This concludes the proof of the real Schur decomposition theorem.

²For the sake of completeness note that in this case 0 and X are, of course, no $(k-1)$ column and row vectors but $(k-2) \times 2$ and $2 \times (k-2)$ matrices instead.

A.3 Validity of the Thouless Vacuum

To see that the Thouless state defines a Bogoliubov vacuum, we have to show that

$$b_n |0\rangle_b^T \stackrel{!}{=} 0 \quad (\text{A.34})$$

for all Bogoliubov annihilation operators b_n . As a preparation, we first write the Thouless state as

$$|0\rangle_b^T = \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] |0\rangle, \quad (\text{A.35})$$

using Einstein notation for better readability. With this, we note that $\exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right]$ fulfils

$$\exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] \exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] = \mathbb{1}_N, \quad (\text{A.36})$$

because even products of $c^\dagger$ operators commute, so we can treat the operator exponentials like ordinary number exponentials. Equation (A.36) allows us to rewrite Eq. (A.34) as

$$b_n |0\rangle_b^T = \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] \exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] b_n \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] |0\rangle \stackrel{!}{=} 0. \quad (\text{A.37})$$

The interesting part of this expression is

$$\exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] b_n \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] = \exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] (U_{kn}^* c_k + V_{kn}^* c_k^\dagger) \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right], \quad (\text{A.38})$$

where we plugged in the definition $b_n = (U_{kn}^* c_k + V_{kn}^* c_k^\dagger)$ from Eq. (5.135). Since the $c_k^\dagger$ commute with the even products of $c^\dagger$ operators, Eq. (A.38) reduces to

$$U_{kn}^* \exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] c_k \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] + V_{kn}^* c_k^\dagger, \quad (\text{A.39})$$

where the non-trivial transformation of the c_k can be computed as

$$\exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] c_k \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] = \sum_{n=0}^{\infty} \frac{1}{n!} \left[\left(-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right), c_k \right]^{(n)} \quad (\text{A.40})$$

using the Hadamard lemma

$$e^A B e^{-A} = \sum_n \frac{1}{n!} [A, B]^{(n)} \quad (\text{A.41})$$

from the Baker-Campbell-Hausdorff formula. Here, we use the notation

$$[A, B]^{(n)} := \underbrace{[A, [A, \dots [A, B] \dots]]}_{n \text{ times}} \quad (\text{A.42})$$

for the n -fold nested commutator between operators A and B . Explicit calculation shows that

$$\begin{aligned} \left[\left(-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right), c_k \right]^{(1)} &= -\frac{1}{2} S_{ij} [c_i^\dagger c_j^\dagger, c_k] \\ &= -\frac{1}{2} S_{ij} (c_i^\dagger c_j^\dagger c_k - c_k c_i^\dagger c_j^\dagger) \\ &= -\frac{1}{2} S_{ij} (c_i^\dagger (\delta_{jk} - c_k c_j^\dagger) - (\delta_{ik} - c_i^\dagger c_k) c_j^\dagger) \\ &= -\frac{1}{2} S_{ij} (c_i^\dagger \delta_{jk} - \delta_{ik} c_j^\dagger) \\ &= -\frac{1}{2} (S_{ik} c_i^\dagger - S_{kj} c_j^\dagger) \\ &\stackrel{(\diamond)}{=} -\frac{1}{2} (S_{ik} c_i^\dagger + S_{ik} c_i^\dagger) \\ &= -S_{ik} c_i^\dagger \end{aligned} \quad (\text{A.43})$$

where we used the skew-symmetry $S_{ki} = -S_{ik}$ and renamed summation indices in $(\diamond)$. With this, we immediately find

$$\begin{aligned} \left[\left(-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right), c_k \right]^{(2)} &= \frac{1}{4} S_{ij} S_{mn} \left[c_i^\dagger c_j^\dagger, \left[c_m^\dagger c_n^\dagger, c_k \right] \right] \\ &= \frac{1}{4} S_{ij} S_{mk} \left[c_i^\dagger c_j^\dagger, c_m^\dagger \right] \\ &= 0, \end{aligned} \quad (\text{A.44})$$

such that

$$\exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] c_k \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] = c_k - S_{ik} c_i^\dagger. \quad (\text{A.45})$$

We can substitute this back into Eq. (A.39) and find

$$\begin{aligned} U_{kn}^* \left(c_k - S_{ik} c_i^\dagger \right) + V_{kn}^* c_k^\dagger &\stackrel{(\diamond)}{=} U_{kn}^* c_k - S_{ik} U_{kn}^* c_i^\dagger + V_{kn}^* c_k^\dagger \\ &= U_{kn}^* c_k - V_{im}^* U_{mk}^{*-1} U_{kn}^* c_i^\dagger + V_{kn}^* c_k^\dagger \\ &= U_{kn}^* c_k - V_{im}^* \delta_{mn} c_i^\dagger + V_{kn}^* c_k^\dagger \\ &\stackrel{(\star)}{=} U_{kn}^* c_k - V_{kn}^* c_k^\dagger + V_{kn}^* c_k^\dagger \\ &= U_{kn}^* c_k, \end{aligned} \quad (\text{A.46})$$

where we plugged in the matrix elements $S_{ik} = V_{im}^* U_{mk}^{*-1}$ of the S matrix in $(\diamond)$ and renamed the summation index $i \rightarrow k$ in $(\star)$. Combined, we therefore have

$$\exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] b_n \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] = U_{kn}^* c_k. \quad (\text{A.47})$$

If we plug this into Eq. (A.37) we finally get

$$\begin{aligned} b_n |0\rangle_b^T &= \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] \exp \left[-\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] b_n \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] |0\rangle \\ &= \exp \left[\frac{1}{2} S_{ij} c_i^\dagger c_j^\dagger \right] U_{kn}^* c_k |0\rangle \\ &= 0, \end{aligned} \quad (\text{A.48})$$

identifying the Thouless state as a valid Bogoliubov vacuum state.

A.4 Skew-Symmetry of the S -Matrix

The proof of the Robledo overlap formula in Eq. (5.202) and the construction of the Thouless state itself – in particular, the calculation done in Eq. (A.43) – make heavy use of the skew-symmetry of S . Here, we show that if S exists, it is always skew-symmetric. To do this, we take $U^\dagger V^* + V^\dagger U^* = 0$ from Eqs. (5.157) and multiply it by $U^{\dagger-1}$ from the left, which yields

$$U^{\dagger-1}U^\dagger V^* + U^{\dagger-1}V^\dagger U^* = 0 \implies V^* = -U^{\dagger-1}V^\dagger U^*. \quad (\text{A.49})$$

Here, we made use of the assumption that S exists, in which case U^{-1} and hence $U^{\dagger-1}$ exist by definition of S . With this we compute

$$\begin{aligned} S_{ij}^\dagger &= S_{ji} \\ &= V_{j\alpha}^* U_{\alpha i}^{*-1} \\ &\stackrel{(\diamond)}{=} \left(-U_{j\beta}^{\dagger-1} V_{\beta\gamma}^\dagger U_{\gamma\alpha}^* \right) U_{\alpha i}^{*-1} \\ &= -U_{j\beta}^{\dagger-1} V_{\beta\gamma}^\dagger \delta_{\gamma i} \\ &= -U_{j\beta}^{\dagger-1} V_{\beta i}^\dagger \\ &\stackrel{(\star)}{=} -V_{i\beta}^* U_{j\beta}^{-1\dagger} \\ &= -V_{i\beta}^* U_{\beta j}^{-1*} \\ &\stackrel{(*)}{=} -V_{i\beta}^* U_{\beta j}^{*-1} \\ &= -S_{ij}, \end{aligned} \quad (\text{A.50})$$

where we plugged in Eq. (A.49) in $(\diamond)$ and used that taking the adjoint and complex conjugate of an invertible matrix M commute with taking its inverse, i.e. that $M^{\dagger-1} = M^{-1\dagger}$ and $M^{*-1} = M^{-1*}$ in $(\star)$ and $(*)$. This shows that S is always skew-symmetric if it exists.

A.5 The Bloch–Messiah Decomposition

The Bloch–Messiah decomposition (BMD) of an $2L \times 2L$ Bogoliubov transformation matrix

$$B = \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix} \quad (\text{A.51})$$

is a factorisation

$$B = \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix} = \begin{pmatrix} C & 0 \\ 0 & C^* \end{pmatrix} \begin{pmatrix} \bar{U} & \bar{V} \\ \bar{V} & \bar{U} \end{pmatrix} \begin{pmatrix} D^\dagger & 0 \\ 0 & D^\tau \end{pmatrix} = C \bar{B} D^\dagger \quad (\text{A.52})$$

of B into a product of block diagonal unitary matrices C and $D^\dagger$, which are defined in terms unitary $L \times L$ matrices C and D , and a real block matrix $\bar{B}$, the blocks of which are (block) diagonal real matrices of the form

$$\bar{U} = \begin{pmatrix} \mathbb{0}_F & & \\ & \bigoplus_{p \in P} u_p \mathbb{1}_2 & \\ & & \mathbb{1}_E \end{pmatrix} \quad \text{and} \quad \bar{V} = \begin{pmatrix} \mathbb{1}_F & & \\ & \bigoplus_{p \in P} i v_p \sigma_y & \\ & & \mathbb{0}_E \end{pmatrix} \quad (\text{A.53})$$

where $u_p^2 + v_p^2 = 1$ and $u_p, v_p \in \mathbb{R}$ for all $p \in P$. Note that we have $F + 2P + E = L$ by construction.³ Getting rid of the particle hole redundancy we may compactly rewrite Eq. (A.52) as

$$\begin{aligned} U &= C \bar{U} D^\dagger \\ V &= C^* \bar{V} D^\dagger. \end{aligned} \quad (\text{A.54})$$

This formulation of the BMD is frequently encountered in literature. Perhaps, its popularity stems from the fact that in this form it is most natural to identify the BMD algorithm as a particle-hole compatible simultaneous singular value decomposition (SVD) of U and V . For those familiar with SVD and its applications in physics, this viewpoint immediately suggests that BMD may be a powerful tool for analysing, partitioning and truncating B to speed up or even enable some kinds of computations involving it.⁴ However, regardless of what SVD veterans may infer about the usefulness of BMD, the understanding of BMD as simultaneous and particle-hole compatible SVDs of U and V provides a powerful construction scheme for the explicit BMD matrices of a given Bogoliubov matrix B . We will now discuss this SVD based construction scheme. As a starting point we have some Bogoliubov matrix B as in Eq. (A.51). First, we use its unitarity

$$B^\dagger B = \begin{pmatrix} U^\dagger & V^\dagger \\ V^\tau & U^\tau \end{pmatrix} \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix} = \begin{pmatrix} \mathbb{1}_L & 0 \\ 0 & \mathbb{1}_L \end{pmatrix} = \begin{pmatrix} U & V^* \\ V & U^* \end{pmatrix} \begin{pmatrix} U^\dagger & V^\dagger \\ V^\tau & U^\tau \end{pmatrix} = B B^\dagger \quad (\text{A.55})$$

to get the conditions

$$\begin{aligned} U^\dagger U + V^\dagger V &= \mathbb{1}_L = U U^\dagger + V^* V^\tau \\ U^\dagger V^* + V^\dagger U^* &= 0 = U V^\dagger + V^* U^\tau \\ V^\tau U + U^\tau V &= 0 = V U^\dagger + U^* V^\tau \\ V^\tau V^* + U^\tau U^* &= \mathbb{1}_L = V V^\dagger + U^* U^\tau. \end{aligned} \quad (\text{A.56})$$

Next we construct the auxiliary matrices $P = U^\dagger U$ and $Q = V^\dagger V$ which are Hermitian and positive semi-definite by construction. Using some of the identities (indicated by parenthesis placement) in Eq. (A.56)

³The exact labels E, P, F of the subscripts anticipate a result about the physical meaning of the corresponding blocks that we will soon obtain.

⁴This is, in fact, why we care about BMD too: we need it to define the Bogoliubov ground state in a trouble-free manner!

we can then show that P and Q commute:

$$\begin{aligned}
& P + Q = \mathbb{1}_L \\
\iff & PQ + QQ = Q \\
\iff & U^\dagger UV^\dagger V + V^\dagger (VV^\dagger) V = V^\dagger V \\
\iff & U^\dagger UV^\dagger V + V^\dagger (\mathbb{1}_L - U^* U^\dagger) V = V^\dagger V \\
\iff & U^\dagger UV^\dagger V - V^\dagger U^* (U^\dagger V) = 0 \\
\iff & U^\dagger UV^\dagger V - V^\dagger U^* (-V^\dagger U) = 0 \\
\iff & U^\dagger UV^\dagger V + V^\dagger (-V U^\dagger) U = 0 \\
\iff & U^\dagger UV^\dagger V - V^\dagger V U^\dagger U = 0 \\
\iff & PQ - QP = 0.
\end{aligned} \tag{A.57}$$

As a consequence, P and Q can be simultaneously diagonalised as

$$P = A D_P A^\dagger \quad \text{and} \quad Q = A D_Q A^\dagger \tag{A.58}$$

with a unitary transformation matrix A and diagonal matrices

$$D_P = \text{diag}(p_1, \dots, p_L) \quad \text{and} \quad D_Q = \text{diag}(q_1, \dots, q_L) \tag{A.59}$$

where $p_i, q_i \in \mathbb{R}_{\geq 0}$ for all $i = 1, \dots, L$. Note that the simultaneous diagonalisation of P and Q combined with $P + Q = \mathbb{1}_L$ tells us that

$$D_P + D_Q = \mathbb{1}_L. \tag{A.60}$$

Next we briefly discuss a practical solution to the simultaneous diagonalisation problem.

Simultaneous Diagonalisation Scheme for Commuting Hermitian Matrices

In practice, the simultaneous diagonalisation matrix A can be obtained as follows. First one diagonalises either P or Q . Say we choose to diagonalise P and find

$$P = B D_P B^\dagger. \tag{A.61}$$

The next step is to compute

$$Q' = B^\dagger Q B. \tag{A.62}$$

Due to the fact that P and Q are simultaneously diagonalisable, we know that Q' is bound to be block diagonal with blocks the size of eigenvalue degeneracies of P . This is because if $\mathbf{v}$ is an eigenvector of P with eigenvalue p then

$$P (Q \mathbf{v}) = Q P \mathbf{v} = Q p \mathbf{v} = p (Q \mathbf{v}), \tag{A.63}$$

tells us that for every eigenvector $\mathbf{v}$ with eigenvalue p the vector $\mathbf{w} := Q \mathbf{v}$ is also an eigenvector of P with eigenvalue p . Therefore, Q maps the eigenspaces of P to themselves. Recall that the columns of B are just the normalised eigenvectors of P with eigenvalues given by the diagonal of D_P so the invariance of the eigenspaces of P with respect to Q means that the worst Q can do to the transformation matrix B is to mix those column eigenvectors that belong to the same eigenvalue. More explicitly, say P has m distinct eigenvalues p_i with degeneracies n_i naturally satisfying $\sum_{i=1}^m n_i = L$ and we write the eigenvalue equations as

$$P \mathbf{v}_{p_i}^j = p_i \mathbf{v}_{p_i}^j \tag{A.64}$$

for $j = 1, \dots, n_i$ and $i = 1, \dots, m$ then

$$QB = Q \begin{pmatrix} v_{p_1,1}^1 & \cdots & v_{p_1,1}^{n_1} & \cdots & v_{p_m,1}^{n_m} \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ v_{p_1,L}^1 & \cdots & v_{p_1,L}^{n_1} & \cdots & v_{p_m,L}^{n_m} \end{pmatrix} = \begin{pmatrix} w_{p_1,1}^1 & \cdots & w_{p_1,1}^{n_1} & \cdots & w_{p_m,1}^{n_m} \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ w_{p_1,L}^1 & \cdots & w_{p_1,L}^{n_1} & \cdots & w_{p_m,L}^{n_m} \end{pmatrix} \tag{A.65}$$

where the column vectors $\mathbf{w}_{p_i}^j$ still satisfy

$$P \mathbf{w}_{p_i}^j = p_i \mathbf{w}_{p_i}^j. \quad (\text{A.66})$$

Since distinct eigenspaces are mutually orthogonal with respect to the inner product, the inner products between $\mathbf{v}$'s and $\mathbf{w}$'s may be spelled out as

$$\mathbf{v}_{p_i}^{j\dagger} \mathbf{w}_{p_k}^l = \delta_{jk} c_{ijl} \quad (\text{A.67})$$

with an i, j, l dependent complex number c_{ijl} and we get

$$\begin{aligned} Q' &= B^\dagger Q B \\ &= \begin{pmatrix} v_{p_1,1}^{1*} & \cdots & v_{p_1,L}^{1*} \\ \vdots & \ddots & \vdots \\ v_{p_1,1}^{n_1*} & \cdots & v_{p_1,L}^{n_1*} \\ \vdots & \ddots & \vdots \\ v_{p_m,1}^{n_m*} & \cdots & v_{p_m,L}^{n_m*} \end{pmatrix} \begin{pmatrix} w_{p_1,1}^1 & \cdots & w_{p_1,1}^{n_1} & \cdots & w_{p_m,1}^{n_m} \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ w_{p_1,L}^1 & \cdots & w_{p_1,L}^{n_1} & \cdots & w_{p_m,L}^{n_m} \end{pmatrix} \\ &= \begin{pmatrix} Q_1 & 0 & \cdots & 0 \\ 0 & Q_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & Q_m \end{pmatrix} \end{aligned} \quad (\text{A.68})$$

with the Q_i being Hermitian $n_i \times n_i$ matrices. In a final step we diagonalise those Q_i that are not yet diagonal, i.e. we find matrices C_i with $i = 1, \dots, m$ such that

$$Q_i = C_i D_{Q_i} C_i^\dagger \quad (\text{A.69})$$

with D_{Q_i} real diagonal matrices. Defining the block diagonal unitary matrix

$$C = \bigoplus_{i=1}^m C_i, \quad (\text{A.70})$$

we can finally define a mutual diagonalisation matrix A of P and Q as

$$A = B C. \quad (\text{A.71})$$

Importantly, the additional application of C does not undo the diagonalisation of P since C only unitarily stirs up the eigenspaces of P . It is worth mentioning that in practice it is the rule rather than the exception that the matrix Q' is readily diagonal so more often than not the only thing that needs to be done in order to find a simultaneous diagonalisation for two commuting Hermitian matrices is to diagonalise either one of them.

The fact that P and Q are positive semi-definite allows us to rewrite D_P and D_Q from Eq. (A.59) as

$$D_P =: \Gamma^2 \quad \text{and} \quad D_Q =: \Lambda^2 \quad (\text{A.72})$$

with

$$\Gamma := \text{diag}(\gamma_1, \dots, \gamma_L) \quad \text{and} \quad \Lambda := \text{diag}(\lambda_1, \dots, \lambda_L) \quad (\text{A.73})$$

where

$$p_i =: \gamma_i^2 \quad \text{and} \quad q_i =: \lambda_i^2 \quad (\text{A.74})$$

and $\gamma_i, \lambda_i \in \mathbb{R}_{>0}$ which combined with Eq. (A.58) promotes SVD's

$$U = W_U \Gamma A^\dagger \quad \text{and} \quad V = W_V \Lambda A^\dagger \quad (\text{A.75})$$

of U and V since

$$\begin{aligned} P &= U^\dagger U = A \Gamma^\dagger W_U^\dagger W_U \Gamma A^\dagger = A \Gamma \Gamma A^\dagger = A \Gamma^2 A^\dagger = A D_P A^\dagger \\ Q &= V^\dagger V = A \Lambda^\dagger W_V^\dagger W_V \Lambda A^\dagger = A \Lambda \Lambda A^\dagger = A \Lambda^2 A^\dagger = A D_Q A^\dagger \end{aligned} \quad (\text{A.76})$$

gives the correct simultaneous diagonalisation due to W_U, W_V unitary and Γ, Λ Hermitian. Now Eq. (A.75) uniquely defines the left unitary transformation matrices

$$W_U = U A \Gamma^{-1} \quad \text{and} \quad W_V = V A \Lambda^{-1} \quad (\text{A.77})$$

whenever Γ and Λ are invertible. The invertibility of Γ (Λ) is lost if (and only if) there exist zero eigenvalues of P (Q) that appear again as zero singular values of U (V). The physical situations in which the BMD creates significant added value are precisely such situations in which V does have singular values of zero and it turns out that such situations occur rather commonly in our models of nature, too. So we are very much interested in creating a BMD algorithm that can function even in the presence of non-invertible Λ and Γ . We can approach a solution to this problem by noting that in the presence of l (m) zero and $k = L - l$ ($n = L - m$) non-zero singular values of U (V) the diagonal matrices Γ and Λ take the form

$$\Gamma = \begin{pmatrix} \mathbf{0}_{l \times l} & \mathbf{0}_{l \times k} \\ \mathbf{0}_{k \times l} & \tilde{\Gamma}_{k \times k} \end{pmatrix} \quad \text{and} \quad \Lambda = \begin{pmatrix} \tilde{\Lambda}_{n \times n} & \mathbf{0}_{n \times m} \\ \mathbf{0}_{m \times n} & \mathbf{0}_{m \times m} \end{pmatrix} \quad (\text{A.78})$$

where we assumed WLOG that the positive semi-definite singular values on the diagonal of Γ are sorted in ascending order while those on the diagonal of Λ are sorted in descending order. For convenience, we denote the total index set by $I = [1, \dots, L]$ and define the invertible index subsets as $\tilde{I} = [l + 1, \dots, L]$ and $\tilde{\tilde{I}} = [1, \dots, n]$. We can then find solutions

$$\tilde{W}_{U, L \times k} = U_{L \times L} \tilde{A}_{L \times k} \tilde{\Gamma}_{k \times k}^{-1} \quad \text{and} \quad \tilde{\tilde{W}}_{V, L \times n} = V_{L \times L} \tilde{\tilde{A}}_{L \times n} \tilde{\tilde{\Lambda}}_{n \times n}^{-1} \quad (\text{A.79})$$

for rectangular left unitary matrices. Here we defined the invertible submatrices $\tilde{\Gamma}_{k \times k} = \Gamma|_{\tilde{I} \times \tilde{I}}$ of Γ and $\tilde{\tilde{\Lambda}}_{n \times n} = \Lambda|_{\tilde{\tilde{I}} \times \tilde{\tilde{I}}}$ of Λ , i.e. those diagonal submatrices accounting for the non-zero singular values, and the corresponding rectangular simultaneous diagonalisation matrices $\tilde{A}_{L \times k} = A|_{I \times \tilde{I}}$ and $\tilde{\tilde{A}}_{L \times n} = A|_{I \times \tilde{\tilde{I}}}$ that are obtained by removing those columns of $A_{L \times L}$ that correspond to zero singular values of Γ and Λ respectively. Note that the restricted invertible index sets $\tilde{I}$ and $\tilde{\tilde{I}}$ do not coincide. Their overlap contains the fully paired indices $I_p = \tilde{I} \cap \tilde{\tilde{I}} = [l + 1, \dots, n]$ that are signified by singular values strictly greater than zero and strictly smaller than one and that will become important later on.⁵ Now we can extend the rectangular left unitary matrices $\tilde{W}_U$ and $\tilde{\tilde{W}}_V$ to quadratic left unitary matrices by adding the respectively required number of l and m appropriate columns them. Of course there is a considerable degree of freedom in this procedure: any collection of l (m) column vectors that constitutes an orthonormal basis of the orthogonal complement of the sub Hilbert space spanned by the k (n) existing column vectors will yield a viable left unitary matrix. Another way to put this is that any orthonormal basis of the nullspaces $\ker(\tilde{W}_U^\dagger)$ and $\ker(\tilde{\tilde{W}}_V^\dagger)$ constitutes a viable extension of the rectangular unitary left matrices. The computation of nullspaces is therefore one systematic way of obtaining such an extension. However, it is also a rather resource-intensive endeavour in practice. Luckily, Eqs. (A.56) allow us to obtain a particularly efficient and convenient solution. To see this, we note that while the right singular vectors $A^\dagger$ of the the SVDs in Eq. (A.75) incorporate the simultaneous diagonalisation of P and Q via Eq. (A.76) the left singular vectors W_U and W_V provide diagonalisations

$$\begin{aligned} U U^\dagger &= W_U \Gamma A^\dagger A \Gamma^\dagger W_U^\dagger = W_U \Gamma \Gamma W_U^\dagger = W_U D_P W_U^\dagger \\ V V^\dagger &= W_V \Lambda A^\dagger A \Lambda^\dagger W_V^\dagger = W_V \Lambda \Lambda W_V^\dagger = W_V D_Q W_V^\dagger \end{aligned} \quad (\text{A.80})$$

⁵As an exception, we included subscripts indicating the matrix dimensions to aid the discussion in this paragraph. We will drop these subscripts asap because they hurt the eyes.

of the Hermitian matrices $UU^\dagger$ and $VV^\dagger$. Then

$$\begin{aligned}
& (UU^\dagger)W_U = W_UD_P \\
\iff & (UU^\dagger)^*W_U^* = W_U^*D_P^* \\
\iff & (U^*U^\dagger)W_U^* = W_U^*D_P^* \\
\iff & (\mathbb{1}_L - VV^\dagger)W_U^* = W_U^*D_P^* \\
\iff & -VV^\dagger W_U^* = W_U^*D_P^* - W_U^* \\
\iff & VV^\dagger W_U^* = W_U^*(\mathbb{1}_L - D_P^*) \\
\iff & VV^\dagger W_U^* = W_U^*(\mathbb{1}_L - D_P^*)
\end{aligned} \tag{A.81}$$

tells us that if $\mathbf{v}$ is an eigenvector of $UU^\dagger$ associated with an eigenvalue $\nu \in [0, 1]$ then its complex conjugate $\mathbf{v}^*$ is an eigenvector of $VV^\dagger$ associated with the eigenvalue $(1 - \nu) \in [0, 1]$. Accordingly, if $\mathbf{v}$ is a left singular vector of U that is associated with a singular value $\sigma \in [0, 1]$ then its complex conjugate $\mathbf{v}^*$ is a left singular vector of V associated with a singular value $(1 - \sigma) \in [0, 1]$. Thus, we may simply choose the complex conjugate of the left singular column vectors associated with unity singular values of U for the missing left singular column vectors associated with zero singular values of V and vice versa, i.e. we define

$$\begin{aligned}
W_{U,L \times L} &:= \left[(\tilde{W}_{V,L \times n})|_{I_{\lambda_i=1}}^* \right]_{L \times l} \star \tilde{W}_{U,L \times k} \\
W_{V,L \times L} &:= \tilde{W}_{V,L \times n} \star \left[(\tilde{W}_{U,L \times k})|_{I_{\gamma_i=1}}^* \right]_{L \times m}
\end{aligned} \tag{A.82}$$

where the restriction onto $I_{\lambda_i=1} = [1, \dots, l]$ and $I_{\gamma_i=1} = [n + 1, \dots, L]$ signifies the restriction on the index subsets associated with the respective unity singular values and where the operator $\star$ simply attaches one $L \times x$ matrix to another $L \times y$ matrix along the second axis, i.e. the column axis. Note that due to $l + k = m + n = L$ the result of this procedure is in fact two $L \times L$ matrices. Furthermore, we have to attach the missing columns for the unitary left singular transformation matrix of U (V) on the left (right) because we assumed the singular values U (V) to be ordered in ascending (descending) order. Now in order for Eq. (A.75) to properly resemble our desired BMD equations in Eq. (A.54) we still have to put in some work. The shared right unitary transformation matrix $A^\dagger$ seems to be neatly in place and the singular value matrices Γ and Λ are already real, although both are still diagonal instead of one being in diagonal and the other in canonical form. We will consider this to be a minor issue. The most serious deviation from the intended BMD form is the lack of a mutual relationship

$$W_U = W_V^* \tag{A.83}$$

between the left unitary matrices. Let us accept the issue of the diagonal/canonical form of Γ and Λ for the time being and concentrate on the more pressing absence of the conjugation relation Eq. (A.83).⁶ To solve this problem, we follow the idea of [DIG UP] and promote Eq. (A.83) to a balancing condition

$$W_U \stackrel{!}{=} W_V^* \tag{A.84}$$

by sticking an exclamation mark on top of the equal sign. We will see that the implementation of this condition will, indeed, involve a necessary manipulation of (either Γ or) Λ that readily puts it into the wanted “canonic” form.⁷ Note that our specific extension strategy in Eq. (A.82) conveniently ensures that Eq. (A.83) is automatically fulfilled everywhere but in the paired index sector $I_p = [l + 1, \dots, n]$ so we only need to work out the paired index sector now. Let us outline the overall balancing strategy in advance to make it easier to follow. The idea is to find a unitary matrix G of the same block diagonal structure as Γ and Λ such that

$$W_U^* = W_V G \tag{A.85}$$

⁶As is good scientific practice, we will secretly hope that the smaller (diagonal/canonic form) problem will resolve itself in the course of solving the bigger, but related (conjugation relation) problem. (It will.)

⁷Following the canonic way of doing things, we will choose to manipulate Λ such that $\bar{V}$ will be the one given in canonic real rather than diagonal real form.

and

$$V = W_V \Lambda A^\dagger = W_V G G^\dagger \Lambda A^\dagger = W_U^* G^\dagger \Lambda A^\dagger = W_U^* \Lambda G^\dagger A^\dagger \quad (\text{A.86})$$

where we used that Λ and G have the same block diagonal structure for the last equality. At this point V has the correct left unitary matrix W_U^* in its BMD but only at the expense of the right unitary matrix changing to $G^\dagger A^\dagger$. To make up for this, we need to find a factorisation of the unusual form

$$G = X K X^\dagger \quad (\text{A.87})$$

with another block diagonal unitary matrix X and some well-behaved auxiliary block diagonal matrix K such that

$$\begin{aligned} V &= W_U^* \Lambda G^\dagger A^\dagger & \text{while} & & U &= W_U \Gamma A^\dagger \\ &= W_U^* \Lambda (X K X^\dagger)^\dagger A^\dagger & & & &= W_U X X^\dagger \Gamma A^\dagger \\ &= W_U^* \Lambda X^* K^\dagger X^\dagger A^\dagger & & & &= (W_U X) \Gamma (X^\dagger A^\dagger) \\ &= (W_U^* X^*) (\Lambda K^\dagger) (X^\dagger A^\dagger) & & & &= C \bar{U} D^\dagger \\ &= C^* \bar{V} D^\dagger \end{aligned} \quad (\text{A.88})$$

where we used that $[\Lambda, X] = [\Gamma, X] = 0$ due to the shared block diagonal form of all these matrices and where we defined

$$C = W_U X \quad \text{and} \quad D = A X \quad (\text{A.89})$$

along with

$$\bar{V} = \Lambda K^\dagger \quad \text{and} \quad \bar{U} = \Gamma. \quad (\text{A.90})$$

Note that the we require the peculiar combination of unitary X and $X^\dagger$ to incorporate the complex conjugation relation. For this to be a useful scheme we would of course need a systematic a way to come up with the required unitary matrix G and its factorisation Eq. (A.87) - so how do we do this in practice? From Eq. (A.85) we may naively guess G as

$$G = W_V^\dagger W_U^*, \quad (\text{A.91})$$

which clearly satisfies Eq. (A.85). For a start this definition of G is readily unitary because

$$\begin{aligned} G G^\dagger &= W_V^\dagger W_U^* (W_V^\dagger W_U^*)^\dagger = W_V^\dagger W_U^* W_U^\dagger W_V = \mathbb{1}_L \\ G^\dagger G &= (W_V^\dagger W_U^*)^\dagger W_V^\dagger W_U^* = W_U^\dagger W_V W_V^\dagger W_U^* = \mathbb{1}_L \end{aligned} \quad (\text{A.92})$$

due to the unitarity of W_U and W_V . Furthermore, we can show that it has the same block diagonal structure as W_V and W_U too. To achieve this, we once again rely on the fact that if $\mathbf{v}$ is a left singular vector of V associated with a singular value $\sigma \in [0, 1]$ of V then $\mathbf{v}^*$ is a left singular vector of U associated with the singular value $\sqrt{1 - \sigma^2} \in [1, 0]$. This tells that even if the σ column vectors of W_V are not generally going to be equal to the complex conjugate $(1 - \sigma)$ column vectors of U they still span the same sub Hilbert space, i.e. if σ is a $2N$ fold degenerate singular value of V with left singular vectors $\mathbf{v}_{\sigma,1}, \dots, \mathbf{v}_{\sigma,2N}$ spanning a sub Hilbert space $\mathcal{H}|_\sigma = \text{span}(\mathbf{v}_{\sigma,1}, \dots, \mathbf{v}_{\sigma,2N})$ and if $\mu = \sqrt{1 - \sigma^2}$ is the associated $2N$ fold degenerate singular value of U with left singular vectors $\mathbf{w}_{\sigma,1}, \dots, \mathbf{w}_{\sigma,2N}$ spanning a sub Hilbert space $\mathcal{H}|_{\mu=\sqrt{1-\sigma^2}} = \text{span}(\mathbf{w}_{\mu=\sqrt{1-\sigma^2},1}, \dots, \mathbf{w}_{\mu=\sqrt{1-\sigma^2},2N})$ then

$$\mathcal{H}|_{\mu=\sqrt{1-\sigma^2}} = \text{span}(\mathbf{w}_{\mu=\sqrt{1-\sigma^2},1}, \dots, \mathbf{w}_{\mu=\sqrt{1-\sigma^2},2N}) = \text{span}(\mathbf{v}_{\sigma,1}^*, \dots, \mathbf{v}_{\sigma,2N}^*) = \mathcal{H}|_\sigma^*. \quad (\text{A.93})$$

We may therefore expand each $\mathbf{w}_{\mu=\sqrt{1-\sigma^2},i}$ in the $\mathbf{v}_i^*$ as

$$\mathbf{w}_{\mu=\sqrt{1-\sigma^2},i} = \sum_{j=1}^{2N} c_{i,j} \mathbf{v}_{\sigma,j}^* \quad (\text{A.94})$$

and vice versa. Thus, we can write the matrix elements of G from Eq. (A.91) as vector products

$$\begin{aligned}
G_{il} &= \mathbf{v}_{\sigma_i, n_{\sigma_i}}^\dagger \mathbf{w}_{\mu_l, m_{\mu_l}}^* \\
&= \mathbf{v}_{\sigma_i, n_{\sigma_i}}^\dagger \mathbf{w}_{\sqrt{1-\sigma_l^2}, m_{\mu_l}}^* \\
&= \mathbf{v}_{\sigma_i, n_{\sigma_i}}^\dagger \left(\sum_{j=1}^{2N_l} c_{m_{\mu_l}, j} \mathbf{v}_{\sigma_l, j}^* \right)^* \\
&= \sum_{j=1}^{2N_l} c_{m_{\mu_l}, j}^* \mathbf{v}_{\sigma_i, n_{\sigma_i}}^\dagger \mathbf{v}_{\sigma_l, j} \\
&= \sum_{j=1}^{2N_l} c_{m_{\mu_l}, j}^* \delta_{\sigma_i, \sigma_l} \delta_{n_{\sigma_i}, j} \\
&= c_{m_{\mu_l}, n_{\sigma_i}}^* \delta_{\sigma_i, \sigma_l}
\end{aligned} \tag{A.95}$$

where we used a multi index notation $i = (\sigma_i, n_{\sigma_i})$ with $\sigma_i \in \{\sigma_1, \dots, \sigma_L\}$ refers to the complete collection of singular values (including degeneracies) while n_{σ_i} labels the distinct left singular vectors of σ_i to keep track of degeneracies. The $\delta_{\sigma_i, \sigma_l}$ in the final line of Eq. (A.95) reveals that G has indeed the same block diagonal structure as W_U and W_V . Note again that due to our construction of W_U and W_V , the matrix G is equal to the identity everywhere but on the blocks corresponding to paired singular values $\sigma, \mu \in (0, 1)$. Finally, without assuming anything other about G than its unitarity we get

$$\begin{aligned}
V &= W_V \Lambda A^\dagger \\
&= W_V G G^\dagger \Lambda A^\dagger \\
&= W_U^* G^\dagger \Lambda A^\dagger
\end{aligned} \tag{A.96}$$

so we have

$$\begin{aligned}
V^\dagger U &= A^* \Lambda G^* W_U^\dagger W_U \Gamma A^\dagger & \text{and} & & U^\dagger V &= A^* \Gamma W_U^\dagger W_U^* G^\dagger \Lambda A^\dagger \\
&= A^* \Lambda G^* \Gamma A^\dagger & & & &= A^* \Gamma G^\dagger \Lambda A^\dagger
\end{aligned} \tag{A.97}$$

which together with $0 = U^\dagger V + V^\dagger U$ from Eqs. (A.56) tells us that

$$\begin{aligned}
0 &= A^* \Gamma G^\dagger \Lambda A^\dagger + A^* \Lambda G^* \Gamma A^\dagger \\
&= \Gamma G^\dagger \Lambda + \Lambda G^* \Gamma \\
&= \Gamma G^\dagger \Lambda + \Lambda G \Gamma.
\end{aligned} \tag{A.98}$$

If we use Einstein summation notation to write this in components we get

$$\begin{aligned}
0 &= \Gamma_{ij} G_{jk}^\dagger \Lambda_{kl} + \Lambda_{ij} G_{jk} \Gamma_{kl} \\
&= \Gamma_i \delta_{ij} G_{jk}^\dagger \delta_{kl} \Lambda_l + \Lambda_i \delta_{ij} G_{jk} \delta_{kl} \Gamma_l \\
&= \Gamma_i G_{il}^\dagger \Lambda_l + \Lambda_i G_{il} \Gamma_l \\
&= \Gamma_i \Lambda_l G_{il}^\dagger + \Gamma_l \Lambda_i G_{il} \\
&\stackrel{(*)}{=} 2D_{il} G_{il}^\dagger + 2D_{li} G_{il} \\
&= 2D_{il} G_{il}^\dagger + 2D_{li} G_{il} + D_{li} G_{il}^\dagger - D_{li} G_{il}^\dagger + D_{il} G_{il} - D_{il} G_{il} \\
&= (D_{il} + D_{li})(G_{il}^\dagger + G_{il}) + (D_{il} - D_{li})(G_{il}^\dagger - G_{il}) \\
&= \square_{il}(G_{il}^\dagger + G_{il}) + \Delta_{il}(G_{il}^\dagger - G_{il})
\end{aligned} \tag{A.99}$$

where we used that Γ and Λ are diagonal, then for convenience multiplied the whole equation by 2 in $(*)$ and finally defined an auxiliary matrix $D_{il} = \Gamma_i \Lambda_l$ along with its symmetric and skew-symmetric parts $\square_{il} = (D_{il} + D_{li}) = (D_{il} + D_{il}^\dagger)$ and $\Delta_{il} = (D_{il} - D_{li}) = (D_{il} - D_{il}^\dagger)$.

Now we differentiate between two cases:

1. $D_{il} = D_{li} =: D$:
 - (a) $D > 0$: $2D(G_{il}^\top + G_{il}) = 0 \implies G_{il}^\top = -G_{il}$
 - (b) $D = 0$: always fulfilled $\implies G_{il}$ arbitrary
2. $D_{il} \neq D_{li}$ where WLOG we assume that $D_{il} > D_{li}$ so $\square_{il} > 0$ and $\Delta_{il} > 0$:

$$\begin{aligned}
 0 &= \square_{il}(G_{il}^\top + G_{il}) + \Delta_{il}(G_{il}^\top - G_{il}) \\
 \implies 0 &= (G_{il}^\top + G_{il}) = (G_{il}^\top - G_{il}) \\
 \implies 0 &= G_{il}^\top = G_{il} = -G_{il}.
 \end{aligned} \tag{A.100}$$

Using the definitions $D_{il} = \Gamma_i \Lambda_l$ and $\Gamma_i^2 + \Lambda_i^2 = 1$ we may rewrite $D_{il} = D_{li}$ as

$$\Gamma_i \Lambda_l = \Gamma_l \Lambda_i \iff \Gamma_i \sqrt{1 - \Gamma_l^2} = \Gamma_l \sqrt{1 - \Gamma_i^2} \iff \frac{\Gamma_i}{\sqrt{1 - \Gamma_i^2}} = \frac{\Gamma_l}{\sqrt{1 - \Gamma_l^2}} \tag{A.101}$$

which for $\Gamma_i, \Gamma_l \in [0, 1)$ only admits the solution⁸

$$\Gamma_i = \Gamma_l. \tag{A.102}$$

With this the 2. case where $D_{il} \neq D_{li}$ just becomes another argument to show that $G_{il} = 0$ whenever $\Gamma_i \neq \Gamma_l$, i.e. that $G = W_V^\dagger W_U^*$ has the same block diagonal structure as W_U and W_V . However, the 1. case where $D_{il} = D_{li} = D$ tells us something new. While for $D = 0$, i.e., for $\Gamma_i, \Lambda_l \in \{0, 1\}$, we only find that G is arbitrary, the $D > 0$ case, i.e., for $\Gamma_i, \Lambda_l \in (0, 1)$, shows that G is necessarily skew-symmetric on the possibly degenerate fully paired blocks. In summary, the matrix G as defined in Eq. (A.91) is unitary, has the same block diagonal structure as W_U and W_V , is equal to the identity on the singular value blocks associated with singular values $\sigma, \mu \in \{0, 1\}$ and is skew-symmetric on the fully paired blocks. We may write

$$G = \mathbb{1}_l \oplus G_1 \oplus \cdots \oplus G_p \oplus \mathbb{1}_m \tag{A.103}$$

where $\mathbb{1}_l$ and $\mathbb{1}_m$ are the unity blocks associated with the zero and one singular values and where the G_i are $2N_i \times 2N_i$ skew-symmetric blocks associated with the N_i -fold degenerate i -th singular value respectively.⁹ Having found a legitimate matrix G we are now in a position to think about how to factorise it in the desired fashion of Eq. (A.87). The unity blocks of G are of course trivial and pose no problem. The skew-symmetric blocks do need some treatment but due to the block diagonal structure of G we can fix them one at a time. Therefore, the problem of factorising all of G at once becomes the much smaller problem of factorising the individual blocks G_i . This is rather useful because the G_i blocks are not only much smaller and easier to handle, but also have the strong symmetric property of being skew-symmetric, which is a significant advantage over G as a whole. For example there is a pretty elegant and well-known solution of the factorisation problem Eq. (A.87) for symmetric matrices which is called the Takagi factorisation. It asserts that any complex symmetric matrix M may be written as

$$M = U D U^\top \tag{A.104}$$

where D is diagonal and where U is unitary. Luckily, there is an analogous version of Takagi's symmetric matrix factorisation that works for skew-symmetric matrices. It is called the Youla decomposition of a skew-symmetric matrix and it states that any complex skew-symmetric matrix S can be factorised as

$$S = U K U^\top \tag{A.105}$$

with a unitary matrix U and a skew-symmetric matrix

$$K = 0 \oplus \cdots \oplus K_1 \oplus \cdots \oplus K_n \tag{A.106}$$

⁸The domain $[0, 1)$ of these functions does of course not cover all possible values of the $\Gamma \in [0, 1]$ but this issue is easily resolved: either one changes from Γ to $\Lambda = \sqrt{1 - \Gamma^2}$ which accounts for $\Gamma = 1$ through $\Lambda = 0$ or one rewrites the functions as Λ_i/Γ_i the domain of which is $(0, 1]$ and which support the same line of argument.

⁹In this notation we therefore have $L = l + m + \sum_{i=1}^p N_i$.

where the K_i are blocks are of the form

$$K_i = \begin{pmatrix} 0 & z_i \\ -z_i & 0 \end{pmatrix} = z_i (i\sigma_y) \quad (\text{A.107})$$

and $z_i \in \mathbb{C}$. Like the Takagi factorisation of symmetric matrices, the Youla decomposition combines a unitary transformation matrix U with its transposed counterpart $U^\top$ rather than its adjoint partner $U^\dagger$ which is why we can absorb the complex phase of the z_i into U and thus make the K_i into real matrices. In our case the skew-symmetric start matrices G_i are not only skew-symmetric but also unitary which means that the Youla decomposition provides us with a factorisation

$$G_i = X_i K_i X_i^\top \quad (\text{A.108})$$

where X_i is unitary and where K_i is a skew-symmetric matrix

$$K_i = K_{i,1} \oplus \cdots \oplus K_{i,2n} \quad (\text{A.109})$$

without zero blocks and with two by two diagonal blocks

$$K_i = \begin{pmatrix} 0 & z_i \\ -z_i & 0 \end{pmatrix} \quad (\text{A.110})$$

where $z_i \in U(1) \subset \mathbb{C}$ such that an absorption of the complex phases into X produces even simpler blocks

$$K_i = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix} = -i\sigma_y. \quad (\text{A.111})$$

These properties are a direct consequence of the G_i 's unitarity which can be seen as follows. Let S be skew-symmetric and unitary then

$$\begin{aligned} \mathbb{1} &= S^\dagger S \\ &= U^* K^\dagger U^\dagger U K U^\top \\ &= U^* K^\dagger K U^\top \\ &= U^* \text{diag}(|z_1|^2, |z_1|^2, \dots, |z_n|^2, |z_n|^2) U^\top \\ &= \text{diag}(|z_1|^2, |z_1|^2, \dots, |z_n|^2, |z_n|^2). \end{aligned} \quad (\text{A.112})$$

The Youla decomposition of G as a whole therefore reads

$$\begin{aligned} G &= \mathbb{1}_l \oplus G_1 \oplus \cdots \oplus G_p \oplus \mathbb{1}_m \\ &= (\mathbb{1}_l \oplus X_1 \oplus \cdots \oplus X_p \oplus \mathbb{1}_m) (\mathbb{1}_l \oplus K_1 \oplus \cdots \oplus K_p \oplus \mathbb{1}_m) (\mathbb{1}_l \oplus X_1^\top \oplus \cdots \oplus X_p^\top \oplus \mathbb{1}_m) \\ &= X K X^\top. \end{aligned} \quad (\text{A.113})$$

If we plug in this K into the definition Eq. (A.90) of $\bar{V}$ that we anticipated earlier we get

$$\begin{aligned} \bar{V} &= \Lambda K^\dagger \\ &= \begin{pmatrix} \mathbb{1}_l & 0 & 0 \\ 0 & \Lambda_P & 0 \\ 0 & 0 & \mathbb{0}_m \end{pmatrix} \begin{pmatrix} \mathbb{1}_l & 0 & 0 \\ 0 & K_P & 0 \\ 0 & 0 & \mathbb{1}_m \end{pmatrix} \\ &= \begin{pmatrix} \mathbb{1}_l & 0 & 0 \\ 0 & \bigoplus_{j=1}^p \Lambda_j \mathbb{1}_2 & 0 \\ 0 & 0 & \mathbb{0}_m \end{pmatrix} \begin{pmatrix} \mathbb{1}_l & 0 & 0 \\ 0 & \bigoplus_{j=1}^p i\sigma_y & 0 \\ 0 & 0 & \mathbb{1}_m \end{pmatrix} \\ &= \begin{pmatrix} \mathbb{1}_l & 0 & 0 \\ 0 & \bigoplus_{j=1}^p i\Lambda_j \sigma_y & 0 \\ 0 & 0 & \mathbb{0}_m \end{pmatrix} \end{aligned} \quad (\text{A.114})$$

which is precisely the canonic form of $\bar{V}$ that is part of the BMD statement. Numerically, the Youla decomposition of the G_i blocks can be efficiently achieved using the *pfapack* package subroutine *pfapack.pfaffian.skew_tridiagonalize* which yields the unitary matrix U and a complex matrix K . The complex phases of the K blocks have to be removed by hand but once that is done we are left with the desired decomposition of G that finally gives us the matrices $C, D, \bar{V}$ and $\bar{U}$ as defined in Eq. (A.89) and Eq. (A.90).

A.6 The Relation Between the Product State and the Thouless State

To prove Eq. (5.155), recall that the product and the Thouless state are only simultenaously defined when both U and V are invertible. According to the BMD algorithm, this is only possible when the Bogoliubov vacuum is fully paired, i.e. if $d_E = d_F = 0$ and $2d_P = N$. In that case, we know that the product state of the Bogoliubov vacuum takes the form

$$\begin{aligned}
|0\rangle_b^P &= \det(D^\dagger) \prod_{p=1}^{N/2} \bar{b}_p \bar{b}_{\bar{p}} |0\rangle \\
&\stackrel{(\diamond)}{=} \det(D^\dagger) \left(\prod_{p=1}^{N/2} v_p \right) \prod_{p=1}^{N/2} \left(u_p + v_p \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right) |0\rangle \\
&\stackrel{(\star)}{=} \det(D^\dagger) \left(\prod_{p=1}^{N/2} v_p u_p \right) \prod_{p=1}^{N/2} \left(1 + \frac{v_p}{u_p} \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right) |0\rangle \\
&\stackrel{(*)}{=} \det(D^\dagger) \left(\prod_{p=1}^{N/2} v_p u_p \right) \prod_{p=1}^{N/2} \exp \left[\frac{v_p}{u_p} \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right] |0\rangle \\
&\stackrel{(\triangle)}{=} \det(D^\dagger) \left(\prod_{p=1}^{N/2} v_p u_p \right) \exp \left[\sum_{p=1}^{N/2} \frac{v_p}{u_p} \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right] |0\rangle, \tag{A.115}
\end{aligned}$$

where we plugged in Eq. (5.189) in $(\diamond)$ and used that $u_p, v_p > 0$ for all p in $(\star)$. After that, we utilised that

$$\left(\bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right)^2 = 0, \tag{A.116}$$

such that

$$\exp \left[\frac{v_p}{u_p} \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right] = 1 + \frac{v_p}{u_p} \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \tag{A.117}$$

in $(*)$. Finally, we made use of the fact that pairs $\bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger$ of creation operators commute such that

$$\exp \left[\frac{v_p}{u_p} \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right] = \exp \left[\sum_{p=1}^{N/2} \frac{v_p}{u_p} \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right] \tag{A.118}$$

in $(\triangle)$. Now, we can use the BMD to show that

$$\begin{aligned}
\exp \left[\frac{1}{2} (\mathbf{c}^\dagger)^\top S \mathbf{c}^\dagger \right] &= \exp \left[\frac{1}{2} (\mathbf{c}^\dagger)^\top V^* U^{*-1} \mathbf{c}^\dagger \right] \\
&\stackrel{(\bullet)}{=} \exp \left[\frac{1}{2} (\mathbf{c}^\dagger)^\top C^\dagger \bar{V} D^* D^\top \bar{U}^{-1} C^* \mathbf{c}^\dagger \right] \\
&\stackrel{(\nabla)}{=} \exp \left[\frac{1}{2} (\bar{\mathbf{c}}^\dagger)^\top \bar{V} \bar{U}^{-1} \bar{\mathbf{c}}^\dagger \right] \\
&\stackrel{(\square)}{=} \exp \left[\sum_{p=1}^{N/2} \frac{v_p}{u_p} \bar{c}_p^\dagger \bar{c}_{\bar{p}}^\dagger \right], \tag{A.119}
\end{aligned}$$

where we plugged in the BMD $U = C^\dagger \bar{U} D$ and $V = C^\top \bar{V} D$ in $(\bullet)$, used that D is unitary so that $D^* D^\top = \mathbb{1}_N$ in (∇) , and finally spelled out the matrix product in elements, using that $\bar{U}_{pk} = \frac{1}{u_p} \delta_{pk}$ and $\bar{V}_{pk} = v_p (\delta_{p\bar{k}} - \delta_{\bar{p}k})$ in $(\square)$. Combined, we get

$$|0\rangle_b^P = \det(D^\dagger) \left(\prod_{p=1}^{N/2} v_p u_p \right) \exp \left[\frac{1}{2} (\mathbf{c}^\dagger)^\top S \mathbf{c}^\dagger \right] |0\rangle. \tag{A.120}$$

The only thing left to show is that

$$\begin{aligned}
\text{Pf}(U^\dagger V^*) &\stackrel{(\diamond)}{=} \text{Pf}(D^\dagger \bar{U} C C^\dagger \bar{V} D^*) \\
&= \text{Pf}(D^\dagger \bar{U} \bar{V} D^*) \\
&\stackrel{(*)}{=} \det(D^*) \text{Pf}(\bar{U} \bar{V}) \\
&\stackrel{(*)}{=} \det(D^*) \prod_{p=1}^{N/2} u_p v_p,
\end{aligned} \tag{A.121}$$

where we once more plugged in the BMD of U and V in $(\diamond)$, and used the Pfaffian identity

$$\text{Pf}(BAB^\top) = \det(B) \text{Pf}(A) \tag{A.122}$$

in $(\star)$. For $(*)$ we used that the Pfaffian of a $2n \times 2n$ skew-symmetric tridiagonal matrix is given as

$$\text{Pf} \begin{pmatrix} 0 & a_1 & 0 & 0 & & \\ -a_1 & 0 & 0 & 0 & & \\ 0 & 0 & 0 & a_1 & & \\ 0 & 0 & -a_1 & 0 & & \\ & & & & \ddots & \\ & & & & & 0 & a_n \\ & & & & & -a_n & 0 \end{pmatrix} = \prod_{p=1}^n a_p \tag{A.123}$$

and that $\bar{U} \bar{V}$ is precisely of that form with $a_p = u_p v_p$. If we plug Eq. (A.121) into Eq. (A.120) we get the desired relation

$$|0\rangle_b^P = \text{Pf}(U^\dagger V^*) \exp \left[\frac{1}{2} (\mathbf{c}^\dagger)^\top S \mathbf{c}^\dagger \right] |0\rangle = \text{Pf}(U^\dagger V^*) |0\rangle_b^T \tag{A.124}$$

we claimed in Eq. (5.155).

A.7 The Robledo Overlap Formula for Product States

In order to prove the Robledo overlap formula for product states, we substitute the relation Eq. (5.155) between the Thouless and the product state into Eq. (5.202). This gives

$$\begin{aligned} {}^P_b \langle 0' | 0 \rangle_b^P &= \text{Pf}(U'^\top V') \text{Pf}(U^\dagger V^*) {}^T_b \langle 0' | 0 \rangle_b^T \\ &= (-1)^{\frac{N(N+1)}{2}} \text{Pf}(U'^\top V') \text{Pf}(U^\dagger V^*) \text{Pf} \begin{pmatrix} S & -\mathbb{1}_N \\ \mathbb{1}_N & -S'^* \end{pmatrix}. \end{aligned} \quad (\text{A.125})$$

To bring this into a more convenient form similar to that of Eq. (5.202) we follow the supplemental material of [93] and write

$$\begin{aligned} \text{Pf} \begin{pmatrix} S & -\mathbb{1}_N \\ \mathbb{1}_N & -S'^* \end{pmatrix} &= \text{Pf} \left(\begin{pmatrix} \mathbb{1}_N & 0 \\ S^{-1} & \mathbb{1}_N \end{pmatrix} \begin{pmatrix} S & 0 \\ 0 & -S'^* + S^{-1} \end{pmatrix} \begin{pmatrix} \mathbb{1}_N & -S^{-1} \\ 0 & \mathbb{1}_N \end{pmatrix} \right) \\ &\stackrel{(\Delta)}{=} \det \begin{pmatrix} \mathbb{1}_N & -S^{-1} \\ 0 & \mathbb{1}_N \end{pmatrix} \text{Pf} \begin{pmatrix} S & 0 \\ 0 & -S'^* + S^{-1} \end{pmatrix} \\ &\stackrel{(\bullet)}{=} \text{Pf}(S) \text{Pf}(-S'^* + S^{-1}), \end{aligned} \quad (\text{A.126})$$

where used the Pfaffian identity Eq. (A.122) together with the skew-symmetry $S^\top = -S$ in (Δ) , and computed the determinant

$$\det \begin{pmatrix} \mathbb{1}_N & 0 \\ S^{-1} & \mathbb{1}_N \end{pmatrix} = \det(\mathbb{1}_N) \det(\mathbb{1}_N) = 1 \quad (\text{A.127})$$

in $(\bullet)$. This allows us to write

$${}^P_b \langle 0' | 0 \rangle_b^P = (-1)^{\frac{N(N+1)}{2}} \text{Pf}(U'^\top V') \text{Pf}(U^\dagger V^*) \text{Pf}(S) \text{Pf}(-S'^* + S^{-1}). \quad (\text{A.128})$$

Next, we use the relation

$$\begin{aligned} \text{Pf}(U'^\top V') &= \text{Pf} \left((V'^{-1} V')^\top U'^\top V' (V'^{-1} V') \right) \\ &= \text{Pf} (V'^\top V' \tau^{-1} U'^\top V' V'^{-1} V') \\ &= \text{Pf} (V'^\top V' \tau^{-1} U'^\top V') \\ &\stackrel{(\diamond)}{=} \det(V') \text{Pf}(V'^\top \tau^{-1} U'^\top) \\ &\stackrel{(\star)}{=} \det(V') \text{Pf}(S'^{\dagger -1}) \\ &\stackrel{(*)}{=} \frac{\det(V')}{\text{Pf}(S'^*)}, \end{aligned} \quad (\text{A.129})$$

where we have used Eq. (A.122) in $(\diamond)$ and

$$S^{\dagger -1} = (V^* U^{*1})^{\dagger -1} = (U^\top V^\top)^{-1} = V^\top \tau^{-1} U^\top \quad (\text{A.130})$$

in $(\star)$. The final equality $(*)$ is achieved using the identities

$$\text{Pf}(-A) = (-1)^n \text{Pf}(A) \quad \text{and} \quad \text{Pf}(A^{-1}) = \frac{(-1)^n}{\text{Pf}(A)} \quad (\text{A.131})$$

for an $n \times n$ skew-symmetric matrix A . Specifically, we rewrite

$$\text{Pf}(S'^{\dagger -1}) = \text{Pf}(-S'^*{}^{-1}) = (-1)^N \text{Pf}(-S'^*{}^{-1}) = \frac{(-1)^{2N}}{\text{Pf}(S'^*)} = \frac{1}{\text{Pf}(S'^*)} \quad (\text{A.132})$$

using the skew-symmetry to substitute $S'^{\dagger} = -S'^*$. We can plug Eq. (A.129) into Eq. (A.128) to get

$$\begin{aligned}
{}_b^P\langle 0'|0\rangle_b^P &= (-1)^{\frac{N(N+1)}{2}} \frac{\det(V')}{\text{Pf}(S'^*)} \frac{\det(V^*)}{\text{Pf}(S)} \text{Pf}(S) \text{Pf}(-S'^* + S^{-1}) \\
&= (-1)^{\frac{N(N+1)}{2}} \det(V') \det(V^*) \frac{1}{\text{Pf}(S'^*)} \text{Pf}(-S'^* + S^{-1}) \\
&\stackrel{(\Delta)}{=} (-1)^{\frac{N(N+1)}{2}} \det(V') \det(V^*) ((-1)^N \text{Pf}(S'^*{}^{-1})) ((-1)^N \text{Pf}(S'^* - S^{-1})) \\
&\stackrel{(\bullet)}{=} (-1)^{\frac{N(N+1)}{2}} \det(V') \det(V^*) \det \begin{pmatrix} \mathbb{1}_N & S'^* \\ 0 & \mathbb{1}_N \end{pmatrix} \text{Pf} \begin{pmatrix} S'^*{}^{-1} & 0 \\ 0 & S'^* - S^{-1} \end{pmatrix} \\
&\stackrel{(\nabla)}{=} (-1)^{\frac{N(N+1)}{2}} \det \begin{pmatrix} V' & 0 \\ 0 & V^* \end{pmatrix} \text{Pf} \left(\begin{pmatrix} \mathbb{1}_N & 0 \\ -S'^* & \mathbb{1}_N \end{pmatrix} \begin{pmatrix} S'^*{}^{-1} & 0 \\ 0 & S'^* - S^{-1} \end{pmatrix} \begin{pmatrix} \mathbb{1}_N & S'^* \\ 0 & \mathbb{1}_N \end{pmatrix} \right) \\
&= (-1)^{\frac{N(N+1)}{2}} \det \begin{pmatrix} V' & 0 \\ 0 & V^* \end{pmatrix} \text{Pf} \begin{pmatrix} S'^*{}^{-1} & \mathbb{1}_N \\ -\mathbb{1}_N & -S^{-1} \end{pmatrix} \\
&\stackrel{(\square)}{=} (-1)^{\frac{N(N+1)}{2}} \text{Pf} \left(\begin{pmatrix} V'^{\tau} & 0 \\ 0 & V^{\dagger} \end{pmatrix} \begin{pmatrix} S'^*{}^{-1} & \mathbb{1}_N \\ -\mathbb{1}_N & -S^{-1} \end{pmatrix} \begin{pmatrix} V' & 0 \\ 0 & V^* \end{pmatrix} \right) \\
&= (-1)^{\frac{N(N+1)}{2}} \text{Pf} \left(\begin{pmatrix} V'^{\tau} & 0 \\ 0 & V^{\dagger} \end{pmatrix} \begin{pmatrix} U'V'^{-1} & \mathbb{1}_N \\ -\mathbb{1}_N & V'U'^{-1} - U^*V^*{}^{-1} \end{pmatrix} \begin{pmatrix} V' & 0 \\ 0 & V^* \end{pmatrix} \right) \\
&= (-1)^{\frac{N(N+1)}{2}} \text{Pf} \begin{pmatrix} V'^{\tau}U' & V'^{\tau}V^* \\ -V^{\dagger}V' & U^{\dagger}V^* \end{pmatrix}, \tag{A.133}
\end{aligned}$$

where we used Eqs. (A.131) in (Δ) , and combined

$$\text{Pf}(S'^*{}^{-1}) \text{Pf}(S'^* - S^{-1}) = \text{Pf} \begin{pmatrix} S'^*{}^{-1} & 0 \\ 0 & S'^* - S^{-1} \end{pmatrix}, \tag{A.134}$$

while inserting a one of the form

$$1 = \det \begin{pmatrix} \mathbb{1}_N & 0 \\ -S'^* & \mathbb{1}_N \end{pmatrix} \tag{A.135}$$

in $(\bullet)$. In (∇) we wrote the product of $\det(V')$ and $\det(V^*)$ as the determinant of a suitable block matrix

$$\det(V') \det(V^*) = \det \begin{pmatrix} V' & 0 \\ 0 & V^* \end{pmatrix}. \tag{A.136}$$

Then, we used Eq. (A.123) to write

$$\det \begin{pmatrix} \mathbb{1}_N & S'^* \\ 0 & \mathbb{1}_N \end{pmatrix} \text{Pf} \begin{pmatrix} S'^*{}^{-1} & 0 \\ 0 & S'^* - S^{-1} \end{pmatrix} = \text{Pf} \left(\begin{pmatrix} \mathbb{1}_N & 0 \\ -S'^* & \mathbb{1}_N \end{pmatrix} \begin{pmatrix} S'^*{}^{-1} & 0 \\ 0 & S'^* - S^{-1} \end{pmatrix} \begin{pmatrix} \mathbb{1}_N & S'^* \\ 0 & \mathbb{1}_N \end{pmatrix} \right) \tag{A.137}$$

in (∇) and

$$\det \begin{pmatrix} V' & S'^* \\ 0 & V^* \end{pmatrix} \text{Pf} \begin{pmatrix} S'^*{}^{-1} & 0 \\ 0 & S'^* - S^{-1} \end{pmatrix} = \text{Pf} \left(\begin{pmatrix} V'^{\tau} & 0 \\ 0 & V^{\dagger} \end{pmatrix} \begin{pmatrix} S'^*{}^{-1} & \mathbb{1}_N \\ -\mathbb{1}_N & -S^{-1} \end{pmatrix} \begin{pmatrix} V' & 0 \\ 0 & V^* \end{pmatrix} \right) \tag{A.138}$$

in $(\square)$. In summary, we arrive at the neatly arranged formula

$${}_b^P\langle 0'|0\rangle_b^P = (-1)^{\frac{N(N+1)}{2}} \text{Pf} \begin{pmatrix} V'^{\tau}U' & V'^{\tau}V^* \\ -V^{\dagger}V' & U^{\dagger}V^* \end{pmatrix} \tag{A.139}$$

for the overlap between two distinct product states that Robledo presented in [90].

A.8 TRS, Fourier Transform and Topology of the Haldane Model

Here, we provide a collection of explicit calculations for Chap. 7.

TRS Transformation of the spinful Haldane Term

The TRS transformation of the spinful Haldane term in Eq. (7.15) transforms as

$$\begin{aligned}
\mathcal{T} t_{\text{MF}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha}} e^{i\xi_{jk}} c_{j\alpha}^{\dagger} c_{k\alpha} \mathcal{T}^{\dagger} &= t_{\text{MF}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha}} e^{-i\xi_{jk}} \mathcal{T} c_{j\alpha}^{\dagger} \mathcal{T}^{\dagger} \mathcal{T} c_{k\alpha} \mathcal{T}^{\dagger} \\
&\stackrel{(\diamond)}{=} t_{\text{MF}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\gamma,\eta}} e^{-i\xi_{jk}} \sigma_y^{\alpha\gamma} \sigma_y^{\alpha\eta} c_{j\gamma}^{\dagger} c_{k\eta} \\
&\stackrel{(\star)}{=} -t_{\text{MF}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\gamma,\eta}} e^{-i\xi_{jk}} (\sigma_y^{\gamma\alpha} \sigma_y^{\alpha\eta}) c_{j\gamma}^{\dagger} c_{k\eta} \\
&\stackrel{(*)}{=} -t_{\text{MF}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\gamma,\eta}} e^{-i\xi_{jk}} \delta_{\gamma\eta} c_{j\gamma}^{\dagger} c_{k\eta} \\
&= -t_{\text{MF}} \sum_{\langle\langle j,k \rangle\rangle_{\gamma}} e^{-i\xi_{jk}} c_{j\gamma}^{\dagger} c_{k\gamma}
\end{aligned} \tag{A.140}$$

where we plugged in Eq. (3.15) in $(\diamond)$, applied the skew-symmetry $\sigma_y^{\dagger} = -\sigma_y$ of the y -Pauli matrix in $(\star)$, and used the relation $\sigma_j \sigma_k = \delta_{jk} + i\epsilon_{jkl} \sigma_l$ for the product of Pauli matrices in $(*)$.

Diagonalisation of the Haldane Model in k -Space

We explicitly diagonalise Eq. (7.15) using the Fourier transform $c_j = 1/\sqrt{L} \sum_{\mathbf{k}} e^{i\mathbf{k}\mathbf{R}} c_{\mathbf{k}}$ of the elementary annihilation and creation operators. Since the spin degree of freedom factors in trivially, we consider only one spin projection of Eq. (7.15), i.e. we transform the original spinless Haldane Hamiltonian

$$H_{\text{H}} = -t_{\text{hop}} \sum_{\langle j,k \rangle} c_j^{\dagger} c_k + V \sum_j \epsilon_j c_j^{\dagger} c_j + t_{\text{MF}} \sum_{\langle\langle j,k \rangle\rangle} e^{i\xi_{jk}} c_j^{\dagger} c_k \equiv H_{\text{Hal}}^{\text{NN}} + H_{\text{Hal}}^{\text{pot}} + H_{\text{Hal}}^{\text{NNN}}. \tag{A.141}$$

Note that the honeycomb lattice structure causes some subtleties in combination with the chiral complex hopping $t_{\text{MF}} e^{i\xi_{jk}}$ of the Haldane model. To deal with this, we formally separate the elementary field operators into sublattice- A operators a_j and sublattice- B operators b_j , writing

$$\begin{aligned}
H_{\text{H}} &= -t_{\text{hop}} \sum_{j=1}^L \sum_{k=1}^3 [a_j^{\dagger} b_{j+\nu_k} + b_{j+\nu_k}^{\dagger} a_j] + V \sum_j [a_j^{\dagger} a_j - b_j^{\dagger} b_j] \\
&\quad + t_{\text{MF}} \sum_{j=1}^L \sum_{k=1}^3 \left[e^{i\xi} [a_{j+\mu_k}^{\dagger} a_j + b_{j+\mu_k}^{\dagger} b_j] + e^{-i\xi} [a_j^{\dagger} a_{j+\mu_k} + b_j^{\dagger} b_{j+\mu_k}] \right],
\end{aligned} \tag{A.142}$$

where $j + \nu_k$ denotes the index of the k -th NN site of j and $j + \mu_k$ labels the index of the k -th NNN site in counterclockwise direction. Specifically, we have

$$\mathbf{R}_{j+\nu_k} = \mathbf{R}_j + \mathbf{a}_k \quad \text{and} \quad \mathbf{R}_{j+\mu_k} = \mathbf{R}_j - (-1)^k \mathbf{b}_k \equiv \mathbf{R}_j + \mathbf{b}'_k, \tag{A.143}$$

with $\mathbf{a}_k$ and $\mathbf{k}$ as given in Eqs. (7.11) and (7.12) respectively. We also introduced the three counterclockwise NNN vectors $\mathbf{b}'_k \equiv -(-1)^k \mathbf{b}_k$ for better readability. Note that $k = 1, 2, 3$ for both NNs and NNNs because the remaining three NNNs lie in clockwise direction and are treated separately in Eq. (A.228). This makes it easier to account for the chiral phase $\xi_{jk} = \pm\xi$ which is positive (negative) if the hopping vector $\Delta \mathbf{R}_{jk} = \mathbf{R}_j - \mathbf{R}_k$. Here, we choose to account for the three clockwise NNNs as illustrated in Fig. A.1. With this choice, the NNN at $\mathbf{R}_{j+\mu_k}$ is located in counterclockwise direction ($\xi_{jk} = +\xi$) if $j \in A$, whereas it is located in clockwise direction ($\xi_{jk} = -\xi$) if $j \in B$. We determine the Fourier

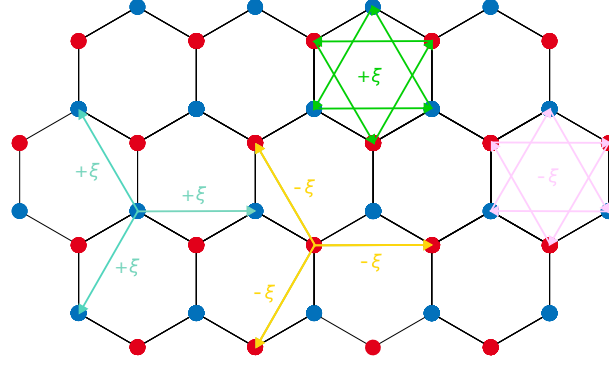


Figure A.1: The blue and red dots label the sites of the two distinct sublattices A and B . Green and pink arrows indicate counterclockwise ($\xi_{jk} = +\xi$) and clockwise ($\xi_{jk} = -\xi$) complex hoppings on both sublattices, while cyan (yellow) arrows mark the three counterclockwise (clockwise) NNN hoppings along $\mathbf{b}_1$, $-\mathbf{b}_2$ and $\mathbf{b}_3$ from a given site in sublattice A (B).

transform of the NN hopping-, sublattice potential- and complex NNN hopping-terms of Eq. (A.228) individually. The NN hopping term yields

$$\begin{aligned}
H_{\text{Hal}}^{\text{NN}} &= -t_{\text{hop}} \sum_{j=1}^L \sum_{k=1}^3 [a_j^\dagger b_{j+\nu_k} + b_{j+\nu_k}^\dagger a_j] \\
&= -t_{\text{hop}} \sum_{j=1}^L \sum_{k=1}^3 \frac{1}{L} \sum_{\mathbf{k}, \mathbf{k}'} [e^{-i\mathbf{k}\mathbf{R}_j} e^{i\mathbf{k}'(\mathbf{R}_j+\mathbf{a}_k)} a_{\mathbf{k}}^\dagger b_{\mathbf{k}'} + e^{-i\mathbf{k}(\mathbf{R}_j+\mathbf{a}_k)} e^{i\mathbf{k}'\mathbf{R}_j} b_{\mathbf{k}}^\dagger a_{\mathbf{k}'}] \\
&= -t_{\text{hop}} \sum_{k=1}^3 \sum_{\mathbf{k}, \mathbf{k}'} \left[\left[\frac{1}{L} \sum_{j=1}^L e^{i(\mathbf{k}'-\mathbf{k})\mathbf{R}_j} \right] e^{i\mathbf{k}'\mathbf{a}_k} a_{\mathbf{k}}^\dagger b_{\mathbf{k}'} + \left[\frac{1}{L} \sum_{j=1}^L e^{i(\mathbf{k}'-\mathbf{k})\mathbf{R}_j} \right] e^{-i\mathbf{k}'\mathbf{a}_k} b_{\mathbf{k}}^\dagger a_{\mathbf{k}'} \right] \\
&\stackrel{(\diamond)}{=} -t_{\text{hop}} \sum_{k=1}^3 \sum_{\mathbf{k}} \left(e^{i\mathbf{k}\mathbf{a}_k} a_{\mathbf{k}}^\dagger b_{\mathbf{k}} + e^{-i\mathbf{k}\mathbf{a}_k} b_{\mathbf{k}}^\dagger a_{\mathbf{k}} \right) \\
&= -t_{\text{hop}} \sum_{k=1}^3 \sum_{\mathbf{k}} \begin{pmatrix} a_{\mathbf{k}}^\dagger & b_{\mathbf{k}}^\dagger \end{pmatrix} \begin{pmatrix} 0 & e^{i\mathbf{k}\mathbf{a}_k} \\ e^{-i\mathbf{k}\mathbf{a}_k} & 0 \end{pmatrix} \begin{pmatrix} a_{\mathbf{k}} \\ b_{\mathbf{k}} \end{pmatrix} \\
&= -t_{\text{hop}} \sum_{k=1}^3 \sum_{\mathbf{k}} \begin{pmatrix} a_{\mathbf{k}}^\dagger & b_{\mathbf{k}}^\dagger \end{pmatrix} \begin{pmatrix} 0 & \cos(\mathbf{k}\mathbf{a}_k) + i\sin(\mathbf{k}\mathbf{a}_k) \\ \cos(\mathbf{k}\mathbf{a}_k) - i\sin(\mathbf{k}\mathbf{a}_k) & 0 \end{pmatrix} \begin{pmatrix} a_{\mathbf{k}} \\ b_{\mathbf{k}} \end{pmatrix} \\
&= -t_{\text{hop}} \sum_{k=1}^3 \sum_{\mathbf{k}} \begin{pmatrix} a_{\mathbf{k}}^\dagger & b_{\mathbf{k}}^\dagger \end{pmatrix} (\cos(\mathbf{k}\mathbf{a}_k)\tau_x - \sin(\mathbf{k}\mathbf{a}_k)\tau_y) \begin{pmatrix} a_{\mathbf{k}} \\ b_{\mathbf{k}} \end{pmatrix}, \tag{A.144}
\end{aligned}$$

while the sublattice potential-term becomes

$$\begin{aligned}
H_{\text{Hal}}^{\text{pot}} &= V \sum_{j=1}^L [a_j^\dagger a_j - b_j^\dagger b_j] \\
&= V \sum_{j=1}^L \frac{1}{L} \sum_{\mathbf{k}, \mathbf{k}'} [e^{-i\mathbf{k}\mathbf{R}_j} e^{i\mathbf{k}'\mathbf{R}_j} a_{\mathbf{k}}^\dagger a_{\mathbf{k}'} - e^{-i\mathbf{k}\mathbf{R}_j} e^{i\mathbf{k}'\mathbf{R}_j} b_{\mathbf{k}}^\dagger b_{\mathbf{k}'}] \\
&= V \sum_{\mathbf{k}, \mathbf{k}'} \left[\left[\frac{1}{L} \sum_{j=1}^L e^{i(\mathbf{k}'-\mathbf{k})\mathbf{R}_j} \right] a_{\mathbf{k}}^\dagger a_{\mathbf{k}'} - \left[\frac{1}{L} \sum_{j=1}^L e^{i(\mathbf{k}'-\mathbf{k})\mathbf{R}_j} \right] b_{\mathbf{k}}^\dagger b_{\mathbf{k}'} \right] \\
&\stackrel{(\diamond)}{=} V \sum_{\mathbf{k}} \begin{pmatrix} a_{\mathbf{k}}^\dagger & a_{\mathbf{k}}^\dagger \\ b_{\mathbf{k}}^\dagger & b_{\mathbf{k}}^\dagger \end{pmatrix} \\
&= V \sum_{\mathbf{k}} \begin{pmatrix} a_{\mathbf{k}}^\dagger & b_{\mathbf{k}}^\dagger \end{pmatrix} \tau_z \begin{pmatrix} a_{\mathbf{k}} \\ b_{\mathbf{k}} \end{pmatrix}, \tag{A.145}
\end{aligned}$$

and the complex NNN hopping-term gives

$$\begin{aligned}
H_{\text{Hal}}^{\text{NNN}} &= t_{\text{MF}} \sum_{j=1}^L \sum_{k=1}^3 \left[e^{i\xi} \left[a_{j+\mu_k}^\dagger a_j + b_j^\dagger b_{j+\mu_k} \right] + e^{-i\xi} \left[a_j^\dagger a_{j+\mu_k} + b_{j+\mu_k}^\dagger b_j \right] \right] \\
&= t_{\text{MF}} \sum_{j=1}^L \sum_{k=1}^3 \frac{1}{L} \sum_{\mathbf{k}, \mathbf{k}'} \left[e^{i\xi} \left[e^{-i\mathbf{k}(\mathbf{R}_j + \mathbf{b}'_k)} e^{i\mathbf{k}'\mathbf{R}_j} a_{\mathbf{k}}^\dagger a_{\mathbf{k}'} + e^{-i\mathbf{k}\mathbf{R}_j} e^{i\mathbf{k}'(\mathbf{R}_j + \mathbf{b}'_k)} b_{\mathbf{k}}^\dagger b_{\mathbf{k}'} \right] \right. \\
&\quad \left. + e^{-i\xi} \left[e^{-i\mathbf{k}\mathbf{R}_j} e^{i\mathbf{k}'(\mathbf{R}_j + \mathbf{b}'_k)} a_{\mathbf{k}}^\dagger a_{\mathbf{k}'} + e^{-i\mathbf{k}(\mathbf{R}_j + \mathbf{b}'_k)} e^{i\mathbf{k}'\mathbf{R}_j} b_{\mathbf{k}}^\dagger b_{\mathbf{k}'} \right] \right] \\
&= t_{\text{MF}} \sum_{k=1}^3 \sum_{\mathbf{k}, \mathbf{k}'} \left[e^{i\xi} \left[\left[\frac{1}{L} \sum_{j=1}^L e^{i(\mathbf{k}' - \mathbf{k})\mathbf{R}_j} \right] e^{-i\mathbf{k}\mathbf{b}'_k} a_{\mathbf{k}}^\dagger a_{\mathbf{k}'} + \left[\frac{1}{L} \sum_{j=1}^L e^{i(\mathbf{k}' - \mathbf{k})\mathbf{R}_j} \right] e^{i\mathbf{k}'\mathbf{b}'_k} b_{\mathbf{k}}^\dagger b_{\mathbf{k}'} \right] \right. \\
&\quad \left. + e^{-i\xi} \left[\left[\frac{1}{L} \sum_{j=1}^L e^{i(\mathbf{k}' - \mathbf{k})\mathbf{R}_j} \right] e^{i\mathbf{k}'\mathbf{b}'_k} a_{\mathbf{k}}^\dagger a_{\mathbf{k}'} + \left[\frac{1}{L} \sum_{j=1}^L e^{i(\mathbf{k}' - \mathbf{k})\mathbf{R}_j} \right] e^{-i\mathbf{k}\mathbf{b}'_k} b_{\mathbf{k}}^\dagger b_{\mathbf{k}'} \right] \right] \\
&\stackrel{(\diamond)}{=} t_{\text{MF}} \sum_{k=1}^3 \sum_{\mathbf{k}} \left[e^{i\xi} \left[e^{-i\mathbf{k}\mathbf{b}'_k} a_{\mathbf{k}}^\dagger a_{\mathbf{k}} + e^{i\mathbf{k}\mathbf{b}'_k} b_{\mathbf{k}}^\dagger b_{\mathbf{k}} \right] \right. \\
&\quad \left. + e^{-i\xi} \left[e^{i\mathbf{k}\mathbf{b}'_k} a_{\mathbf{k}}^\dagger a_{\mathbf{k}} + e^{-i\mathbf{k}\mathbf{b}'_k} b_{\mathbf{k}}^\dagger b_{\mathbf{k}} \right] \right] \\
&= t_{\text{MF}} \sum_{k=1}^3 \sum_{\mathbf{k}} \left[\left[e^{i(\xi - \mathbf{k}\mathbf{b}'_k)} + e^{-i(\xi - \mathbf{k}\mathbf{b}'_k)} \right] a_{\mathbf{k}}^\dagger a_{\mathbf{k}} + \right. \\
&\quad \left. + \left[e^{i(\xi + \mathbf{k}\mathbf{b}'_k)} + e^{-i(\xi + \mathbf{k}\mathbf{b}'_k)} \right] b_{\mathbf{k}}^\dagger b_{\mathbf{k}} \right] \\
&= 2t_{\text{MF}} \sum_{k=1}^3 \sum_{\mathbf{k}} \left[\cos(\xi - \mathbf{k}\mathbf{b}'_k) a_{\mathbf{k}}^\dagger a_{\mathbf{k}} + \cos(\xi + \mathbf{k}\mathbf{b}'_k) b_{\mathbf{k}}^\dagger b_{\mathbf{k}} \right] \\
&= 2t_{\text{MF}} \sum_{k=1}^3 \sum_{\mathbf{k}} \left[\left[\cos(\xi) \cos(-\mathbf{k}\mathbf{b}'_k) - \sin(\xi) \sin(-\mathbf{k}\mathbf{b}'_k) \right] a_{\mathbf{k}}^\dagger a_{\mathbf{k}} + \left[\cos(\xi) \cos(\mathbf{k}\mathbf{b}'_k) - \sin(\xi) \sin(\mathbf{k}\mathbf{b}'_k) \right] b_{\mathbf{k}}^\dagger b_{\mathbf{k}} \right] \\
&= 2t_{\text{MF}} \sum_{k=1}^3 \sum_{\mathbf{k}} \left[\left[\cos(\xi) \cos(\mathbf{k}\mathbf{b}'_k) + \sin(\xi) \sin(\mathbf{k}\mathbf{b}'_k) \right] a_{\mathbf{k}}^\dagger a_{\mathbf{k}} + \left[\cos(\xi) \cos(\mathbf{k}\mathbf{b}'_k) - \sin(\xi) \sin(\mathbf{k}\mathbf{b}'_k) \right] b_{\mathbf{k}}^\dagger b_{\mathbf{k}} \right] \\
&= 2t_{\text{MF}} \sum_{k=1}^3 \sum_{\mathbf{k}} (a_{\mathbf{k}}^\dagger \ b_{\mathbf{k}}^\dagger) \begin{pmatrix} \cos(\xi) \cos(\mathbf{k}\mathbf{b}'_k) + \sin(\xi) \sin(\mathbf{k}\mathbf{b}'_k) & 0 \\ 0 & \cos(\xi) \cos(\mathbf{k}\mathbf{b}'_k) - \sin(\xi) \sin(\mathbf{k}\mathbf{b}'_k) \end{pmatrix} \begin{pmatrix} a_{\mathbf{k}} \\ b_{\mathbf{k}} \end{pmatrix} \\
&= 2t_{\text{MF}} \sum_{k=1}^3 \sum_{\mathbf{k}} (a_{\mathbf{k}}^\dagger \ b_{\mathbf{k}}^\dagger) (\cos(\xi) \cos(\mathbf{k}\mathbf{b}'_k) \mathbb{1}_2 + \sin(\xi) \sin(\mathbf{k}\mathbf{b}'_k) \tau_z) \begin{pmatrix} a_{\mathbf{k}} \\ b_{\mathbf{k}} \end{pmatrix}. \tag{A.146}
\end{aligned}$$

Here we denoted the Pauli matrices on sublattice pseudospin by τ_j to distinguish them from the Pauli matrices σ_j describing electron spin. Combined, the total Fourier transformed spinless Haldane Hamiltonian takes the form

$$H_{\text{Hal}} = \sum_{\mathbf{k}} (a_{\mathbf{k}}^\dagger \ b_{\mathbf{k}}^\dagger) (h_0(\mathbf{k}) \mathbb{1}_2 + \mathbf{h}(\mathbf{k}) \boldsymbol{\tau}) \begin{pmatrix} a_{\mathbf{k}} \\ b_{\mathbf{k}} \end{pmatrix} \tag{A.147}$$

where

$$h_0(\mathbf{k}) = 2t_{\text{MF}} \sum_{j=1}^3 \cos(\xi) \cos(\mathbf{k}\mathbf{b}'_j) \tag{A.148}$$

and

$$h_x(\mathbf{k}) = -t_{\text{hop}} \sum_{j=1}^3 \cos(\mathbf{k}\mathbf{a}_j) \quad , \quad h_y(\mathbf{k}) = t_{\text{hop}} \sum_{j=1}^3 \sin(\mathbf{k}\mathbf{a}_j) \quad , \quad h_z(\mathbf{k}) = V + 2t_{\text{MF}} \sum_{j=1}^3 \sin(\xi) \sin(\mathbf{k}\mathbf{b}'_j). \tag{A.149}$$

We can further simplify this expression by defining

$$x \equiv \frac{\sqrt{3}a_0 k_x}{2} \quad \text{and} \quad y \equiv \frac{a_0 k_y}{2} \quad (\text{A.150})$$

with which

$$\mathbf{k}\mathbf{a}_1 = 2y \quad , \quad \mathbf{k}\mathbf{a}_2 = x - y \quad , \quad \mathbf{k}\mathbf{a}_3 = -x - y \quad (\text{A.151})$$

and

$$\mathbf{k}\mathbf{b}'_1 = \mathbf{k}\mathbf{b}_1 = -x + 3y \quad , \quad \mathbf{k}\mathbf{b}'_2 = -\mathbf{k}\mathbf{b}_2 = -x - 3y \quad , \quad \mathbf{k}\mathbf{a}'_3 = \mathbf{k}\mathbf{b}_3 = 2x \quad (\text{A.152})$$

so that

$$\begin{aligned} h_0(\mathbf{k}) &= 2t_{\text{MF}} \cos(\xi) \sum_{j=1}^3 \cos(\mathbf{k}\mathbf{b}'_j) \\ &= 2t_{\text{MF}} \cos(\xi) \left[\cos(-x + 3y) + \cos(-x - 3y) + \cos(2x) \right] \\ &= 2t_{\text{MF}} \cos(\xi) \left[\cos(-x) \cos(3y) - \sin(-x) \sin(3y) + \cos(-x) \cos(-3y) - \sin(-x) \sin(-3y) + \cos(2x) \right] \\ &= 2t_{\text{MF}} \cos(\xi) \left[\cos(x) \cos(3y) + \cancel{\sin(x) \sin(3y)} + \cos(x) \cos(3y) - \cancel{\sin(x) \sin(3y)} + \cos(2x) \right] \\ &= 2t_{\text{MF}} \cos(\xi) \left[2 \cos(x) \cos(3y) + \cos(2x) \right] \end{aligned} \quad (\text{A.153})$$

and

$$\begin{aligned} h_x(\mathbf{k}) &= -t_{\text{hop}} \sum_{j=1}^3 \cos(\mathbf{k}\mathbf{a}_j) \\ &= -t_{\text{hop}} \left[\cos(2y) + \cos(x - y) + \cos(-x - y) \right] \\ &= -t_{\text{hop}} \left[\cos(2y) + \cos(x) \cos(-y) - \sin(x) \sin(-y) + \cos(-x) \cos(-y) - \sin(-x) \sin(-y) \right] \\ &= -t_{\text{hop}} \left[\cos(2y) + \cos(x) \cos(y) + \cancel{\sin(x) \sin(y)} + \cos(x) \cos(y) - \cancel{\sin(x) \sin(y)} \right] \\ &= -t_{\text{hop}} \left[2 \cos(x) \cos(y) + \cos(2y) \right] \end{aligned} \quad (\text{A.154})$$

$$\begin{aligned} h_y(\mathbf{k}) &= t_{\text{hop}} \sum_{j=1}^3 \sin(\mathbf{k}\mathbf{a}_j) \\ &= t_{\text{hop}} \left[\sin(2y) + \sin(x - y) + \sin(-x - y) \right] \\ &= t_{\text{hop}} \left[\sin(2y) + \sin(x) \cos(-y) + \cos(x) \sin(-y) + \sin(-x) \cos(-y) + \cos(-x) \sin(-y) \right] \\ &= t_{\text{hop}} \left[\sin(2y) + \cancel{\sin(x) \cos(y)} - \cos(x) \sin(y) - \cancel{\sin(x) \cos(y)} - \cos(x) \sin(y) \right] \\ &= -t_{\text{hop}} \left[2 \cos(x) \sin(y) - \sin(2y) \right] \end{aligned} \quad (\text{A.155})$$

$$\begin{aligned}
h_z(\mathbf{k}) &= V + 2t_{\text{MF}} \sin(\xi) \sum_{j=1}^3 \sin(\mathbf{k}\mathbf{b}'_j) \\
&= V + 2t_{\text{MF}} \sin(\xi) \left[\sin(-x + 3y) + \sin(-x - 3y) + \sin(2x) \right] \\
&= V + 2t_{\text{MF}} \sin(\xi) \left[\sin(-x + 3y) - \sin(x + 3y) + \sin(2x) \right] \\
&= V + 2t_{\text{MF}} \sin(\xi) \left[\sin(-x) \cos(3y) + \cos(-x) \sin(3y) - \sin(x) \cos(3y) - \cos(x) \sin(3y) + \sin(2x) \right] \\
&= V + 2t_{\text{MF}} \sin(\xi) \left[-\sin(x) \cos(3y) + \underline{\cos(x) \sin(3y)} - \sin(x) \cos(3y) - \underline{\cos(x) \sin(3y)} + \sin(2x) \right] \\
&= V - 2t_{\text{MF}} \sin(\xi) \left[2 \sin(x) \cos(3y) - \sin(2x) \right].
\end{aligned} \tag{A.156}$$

Chern Number of the Haldane Model

The Haldane model is characterised by the first Chern number, which, as pointed out in Sec. 2.3, obstructs nowhere vanishing global sections. Here we will utilise the (im)possibility of finding nowhere vanishing global sections to compute the Chern number explicitly. The diagonalisation of Eq. (7.17) yields the two valence (−) and conduction (+) eigenstates

$$u_{\pm}(\mathbf{k}) = \frac{1}{N_{\pm}} \begin{pmatrix} h_z(\mathbf{k}) \pm |\mathbf{h}(\mathbf{k})| \\ h_x(\mathbf{k}) + ih_y(\mathbf{k}) \end{pmatrix}, \tag{A.157}$$

where $N_{\pm}$ is the normalising factor and $|\mathbf{h}(\mathbf{k})| = \sqrt{h_x(\mathbf{k})^2 + h_y(\mathbf{k})^2 + h_z(\mathbf{k})^2}$. These correspond to (local) sections of the valence and conduction subbundles. Note that these sections are only locally defined because we find

$$u_{\pm}(\mathbf{K}_{\eta}) = \frac{1}{N_{\pm}} \begin{pmatrix} h_z(\mathbf{K}_{\eta}) \pm |h_z(\mathbf{K}_{\eta})| \\ 0 \end{pmatrix} =: \frac{1}{N_{\pm}} \begin{pmatrix} f_{\pm}(\mathbf{K}_{\eta}) \\ 0 \end{pmatrix} \tag{A.158}$$

the Dirac points $\mathbf{K}_{\eta}$, which vanishes as soon as

$$f_{\pm}(\mathbf{K}_{\eta}) = h_z(\mathbf{K}_{\eta}) \pm |h_z(\mathbf{K}_{\eta})| = 0. \tag{A.159}$$

In the conduction bundle, $f_{\pm}(\mathbf{K}_{\eta})$ vanishes for $h_z(\mathbf{K}_{\eta}) < 0$. In the valence bundle it vanishes for $h_z(\mathbf{K}_{\eta}) > 0$. Exploiting the gauge freedom of the states, we construct an alternative, phase-shifted eigenstate $u_{\pm}^2(\mathbf{k})$ that avoids vanishing at the critical points where the original section fails. The original eigenstate will be denoted by $u_{\pm}^1(\mathbf{k})$ from now on. Each valid gauge, including $u_{\pm}^2(\mathbf{k})$, corresponds to another section of the subbundle. We now suggestively choose the following U(1) gauge transformation:

$$e^{i\varphi_{\pm}(\mathbf{k})} := \frac{\frac{h_z(\mathbf{k}) \mp |\mathbf{h}(\mathbf{k})|}{h_x(\mathbf{k}) + ih_y(\mathbf{k})}}{\left| \frac{h_z(\mathbf{k}) \mp |\mathbf{h}(\mathbf{k})|}{h_x(\mathbf{k}) + ih_y(\mathbf{k})} \right|}. \tag{A.160}$$

With this, the alternative eigenstate becomes

$$u_{\pm}^2(\mathbf{k}) = e^{i\varphi_{\pm}(\mathbf{k})} \cdot u_{\pm}^1(\mathbf{k}) = \frac{1}{N_{\pm}^2} \begin{pmatrix} -\frac{h_x(\mathbf{k})^2 + h_y(\mathbf{k})^2}{h_x(\mathbf{k}) + ih_y(\mathbf{k})} \\ h_z(\mathbf{k}) \mp |h_z(\mathbf{k})| \end{pmatrix} \tag{A.161}$$

with an according normalising factor $N_{\pm}^2$. At the Dirac points this simplifies to

$$u_{\pm}^2(\mathbf{K}_{\eta}) = \frac{1}{N_{\pm}^2} \begin{pmatrix} 0 \\ h_z(\mathbf{K}_{\eta}) \mp |h_z(\mathbf{K}_{\eta})| \end{pmatrix}. \tag{A.162}$$

By construction, $u_{\pm}^2(\mathbf{K}_{\eta})$ displays the opposite behaviour of $u_{\pm}^1(\mathbf{K}_{\eta})$: in the conduction band $u_{+}^2(\mathbf{K}_{\eta})$ vanishes when $h_z(\mathbf{K}_{\eta}) > 0$ while in the valence band $u_{-}^2(\mathbf{K}_{\eta})$ it vanishes for $h_z(\mathbf{K}_{\eta}) < 0$. Since the value

parameters	$h_z(\mathbf{K}_\tau)$	$u_+(\mathbf{k})$	$u_-(\mathbf{k})$
$V > 3\sqrt{3}t_{\text{MF}} \sin(\xi)$	> 0	$u_+^1(\mathbf{k})$	$u_-^2(\mathbf{k})$
$V < -3\sqrt{3}t_{\text{MF}} \sin(\xi)$	< 0	$u_+^2(\mathbf{k})$	$u_-^1(\mathbf{k})$
$-3\sqrt{3}t_{\text{MF}} \sin(\xi) < V < 3\sqrt{3}t_{\text{MF}} \sin(\xi)$	$?$	$?$	$?$

Table A.2: Global phase choices of $u_\pm(\mathbf{k})$ for different parameter regions.

and most notably the sign of $h_z(\mathbf{K}_\eta)$ depend solely on the parameters V and t_{MF} , we may perform a case analysis in their respect. Using $h_z(\mathbf{k})$ from Tab. 7.1 we can easily identify three cases; these are listed in Tab. A.2. In the first two cases the respective choice is unique and good for the entire Brillouin torus. The only configuration of parameters that does not allow for a unique phase choice is given in the third and last case. For $\xi \in (0, \pi]$ we have

$$\left. \begin{array}{l} h_z(\mathbf{K}_{\eta=+1}) > 0 \\ h_z(\mathbf{K}_{\eta=-1}) < 0 \end{array} \right\} \quad \text{if} \quad -3\sqrt{3}t_{\text{MF}} \sin(\xi) < V < 3\sqrt{3}t_{\text{MF}} \sin(\xi), \quad (\text{A.163})$$

which is reversed for $\xi \in (-\pi, 0]$ as then the sign of $\sin(\xi)$ is negative. In either parameter range it is impossible to put down a frame of the subbundle by globally assigning a phase to the state vectors. We have to cut the Brillouin zone into two disjoint parts K^I and K^{II} such that $\mathbf{K}_{\tau=+1} \in K^I$ and $\mathbf{K}_{\tau=-1} \in K^{II}$. Then we define sections $u^I(\mathbf{k})$ on K^I and $u^{II}(\mathbf{k})$ on K^{II} . For $\xi > 0$ we assign the state vectors according to regions as follows

$$u_\pm^I(\mathbf{k}) = \begin{cases} u_+^1(\mathbf{k}) \\ u_-^2(\mathbf{k}) \end{cases} \quad \text{and} \quad u_\pm^{II}(\mathbf{k}) = \begin{cases} u_+^2(\mathbf{k}) \\ u_-^1(\mathbf{k}) \end{cases}. \quad (\text{A.164})$$

For $\xi < 0$ the assignment is reversed. Note that the valence and conduction sections are oppositely defined on both regions. While the conduction bundle is cut in the fashion of $u^1(\mathbf{k})$ in region I , the valence bundle takes the form of $u^2(\mathbf{k})$ there. The converse situation arises on region II . Our I and II versions of the wave function give rise to a Berry connection each:

$$\mathcal{A}_\pm^s(\mathbf{k}) = -i \langle u_\pm^s(\mathbf{k}) | \partial_{\mathbf{k}} | u_\pm^s(\mathbf{k}) \rangle \quad \text{with} \quad s \in \{I, II\}. \quad (\text{A.165})$$

We observe that $u^I(\mathbf{k})$ and $u^{II}(\mathbf{k})$ are related by a ‘‘gauge transformation’’ that is defined according to Eq. (A.161). Due to the inverse assignment of $u^{1/2}(\mathbf{k})$ to the conduction and valence bundle on both regions, however, their phases differ by a sign. Their relation is inherited by the Berry connection as well so we end up with

$$u_\pm^{II}(\mathbf{k}) := e^{\pm i\varphi_\pm(\mathbf{k})} \cdot u_\pm^I(\mathbf{k}) \quad , \quad \mathcal{A}_\pm^{II}(\mathbf{k}) := \mathcal{A}_\pm^I(\mathbf{k}) \pm \partial_{\mathbf{k}} \varphi_\pm(\mathbf{k}). \quad (\text{A.166})$$

Remember that, other than the Berry connection, the Berry curvature $\mathcal{F}_\pm(\mathbf{k})$ is not gauge dependent. It is given by

$$\mathcal{F}_\pm^I(\mathbf{k}) = \mathcal{F}_\pm^{II}(\mathbf{k}) =: \mathcal{F}_\pm(\mathbf{k}). \quad (\text{A.167})$$

Using Eqs. (A.167) and (A.166) in Eq. (2.237) we arrive at a straightforward expression for the Chern number:

$$\begin{aligned} C_\pm &= -\frac{1}{2\pi} \int_{\mathbb{T}^2} dS(\mathbf{k}) \mathcal{F}_\pm(\mathbf{k}) \\ &\stackrel{(\diamond)}{=} -\frac{1}{2\pi} \left[\int_{\partial K^I} d\mathbf{k} \mathcal{A}_\pm^I(\mathbf{k}) + \int_{\partial K^{II}} d\mathbf{k} \mathcal{A}_\pm^{II}(\mathbf{k}) \right] \\ &\stackrel{(*)}{=} -\frac{1}{2\pi} \int_{\partial K^I} d\mathbf{k} (\mathcal{A}_\pm^I(\mathbf{k}) - \mathcal{A}_\pm^{II}(\mathbf{k})) \\ &= \pm \frac{1}{2\pi} \int_{\partial K^I} d\mathbf{k} \partial_{\mathbf{k}} \varphi_\pm(\mathbf{k}). \end{aligned} \quad (\text{A.168})$$

For $(\diamond)$ we split the integral over $\mathbb{T}^2$ into two integrals over $K^{I/II}$ and used Stokes' theorem to recover the boundary integrals over the respective Berry connections. In $(\star)$ we utilised that the region boundaries ∂K^I and ∂K^{II} coincide on $\mathbb{T}^2$ and differ only in their opposite orientation, i.e. $\partial K^I = -\partial K^{II}$.

We seek to parameterise ∂K^I as simply as possible. One practical choice is to set up ∂K^I to be a small circle around $\mathbf{K}_{\tau=+1}$ such that we can Taylor expand h_0 and $\mathbf{h}$ for small deviations $\delta\mathbf{k}$ from $\mathbf{K}_{\tau=+1}$:

$$h_0(\delta\mathbf{k}) \approx -3t_{\text{MF}} \cos(\xi), \quad h_x(\delta\mathbf{k}) \approx t_1 \frac{3a}{2} \delta k_x, \quad (\text{A.169})$$

$$h_y(\delta\mathbf{k}) \approx t_1 \frac{3a}{2} \delta k_y, \quad h_z(\delta\mathbf{k}) \approx V + \tau 3\sqrt{3}t_{\text{MF}} \sin(\xi). \quad (\text{A.170})$$

It is natural to parameterise our circular ∂K^I in polar coordinates. Accordingly, it will come in handy to have a polar coordinate parameterisation of the integrand $\varphi_{\pm}(\mathbf{k})$ at our disposal. This may be obtained following Eq. (A.170) and Eq. (A.161):

$$\begin{aligned} e^{i\varphi_{\pm}(\delta\mathbf{k})} &= \frac{h_z(\delta\mathbf{k}) \mp |\mathbf{h}(\delta\mathbf{k})|}{|h_z(\delta\mathbf{k}) \mp |\mathbf{h}(\delta\mathbf{k})||} \cdot \frac{|h_x(\delta\mathbf{k}) + ih_y(\delta\mathbf{k})|}{h_x(\delta\mathbf{k}) + ih_y(\delta\mathbf{k})} \\ &=: \frac{R(\delta\mathbf{k})}{|R(\delta\mathbf{k})|} \cdot \frac{|\delta k_x + i\delta k_y|}{\delta k_x + i\delta k_y} \\ &= \text{sign}(R(\delta\mathbf{k})) \cdot \frac{|\delta\mathbf{k}|e^{i\theta}}{|\delta\mathbf{k}|e^{i\theta}} \\ &= \text{sign}(R(\delta\mathbf{k})) \cdot e^{-i\vartheta}. \end{aligned} \quad (\text{A.171})$$

The second to last equality formally adapts polar coordinates in the complex $\mathbf{k}$ -plane. Also, we defined $R(\delta\mathbf{k}) := h_z(\delta\mathbf{k}) \mp |\mathbf{h}(\delta\mathbf{k})|$ and utilised that it is real such that $\frac{R(\delta\mathbf{k})}{|R(\delta\mathbf{k})|}$ is the sign of $R(\delta\mathbf{k})$. Obviously, $|\mathbf{h}(\delta\mathbf{k})| \geq |h_z(\delta\mathbf{k})|$. The parameters V and t_{MF} are bigger than or equal to zero such that the sign of $h_z(\delta\mathbf{k})$ is determined by the values of the valley index τ and angle ξ . For now, we shall discuss $\tau = +1$ and $\xi > 0$. This choice makes $h_z(\delta\mathbf{k})$ positive and a subsequent closer look reveals that the sign of $R(\delta\mathbf{k})$ only depends on whether the phase belongs to the valence band or to the conduction band.

At this point we may absorb the respective sign in the form of a constant, band-dependent term added to the exponent:

$$\exp\{i\varphi_{\pm}(\delta\mathbf{k})\} = \exp\left\{i\left(-\vartheta + \frac{(1 \pm 1)}{2}\pi\right)\right\}. \quad (\text{A.172})$$

It is easily confirmed that this construction does, in fact, make up for the sign, depending on whether the valence $(-)$ or conduction $(+)$ band is considered. Most importantly it allows us to compare exponents and find that

$$\varphi_{\pm}(\delta\mathbf{k}) = -\theta + X_{\pm} \quad \text{with } \theta \in [0, 2\pi), \quad (\text{A.173})$$

is independent of $\delta\mathbf{k}$. Note that we defined a band dependent constant $X_{\pm} = \frac{(1 \pm 1)}{2}\pi$. Since we designed ∂K to be a small circle, we employ polar coordinates and find

$$\begin{aligned} C_{\pm} &= \pm \frac{1}{2\pi} \int_{\partial K} d\delta\mathbf{k} \partial_{\mathbf{k}} \varphi_{\pm}(\delta\mathbf{k}) \\ &= \pm \frac{1}{2\pi} \int_0^{2\pi} r' d\theta \left(\partial_r + \frac{1}{r} \partial_{\theta} \right) \varphi_{\pm}(\theta) \\ &= \pm \frac{1}{2\pi} \int_0^{2\pi} d\theta \partial_{\theta} \varphi_{\pm}(\theta) \\ &= \pm \frac{1}{2\pi} [\varphi_{\pm}(\theta)]_0^{2\pi} \\ &= \pm \frac{1}{2\pi} [-\theta \pm X_{\pm}]_0^{2\pi} \\ &= \mp 1. \end{aligned} \quad (\text{A.174})$$

In the second line we transformed to polar coordinates and explicitly parametrised the small circular boundary ∂K . In the next-to-last line we plugged in Eq. (A.173) and used that the constant term vanishes due to definite integration.

The above computation was done for $\tau = +1$ and $\xi > 0$. We saw that for $\xi < 0$ the choice of state vectors is reversed such that, for fixed $\tau = +1$, this change gives a global sign for $\varphi_{\pm}(\delta\mathbf{k})$ and hence for C . The choice of $\tau = +1$ in itself is a more technical issue as it is really only a matter of computing the invariant. We have chosen to integrate around the $\mathbf{K}(\tau = +1)$ point but we could have performed the calculation around $\mathbf{K}'(\tau = -1)$ just as well. Choosing $\tau = -1$ instead of $\tau = +1$ reverses the direction of the oriented boundary between K^I and K^{II} which results in a global sign for C . This sign, however, is compensated by the simultaneously reversed phase convention for the valence and conduction states. Importantly, the above computation shows that for any fixed configuration of parameters, the Chern numbers of the conduction and valence bundle differ by a sign, that is

$$C_+ = -C_- \quad \Longleftrightarrow \quad C_+ + C_- = 0. \quad (\text{A.175})$$

This result is consistent with an important fact that we have mentioned before: the total Bloch bundle is always bound to be trivial.

A.9 The Magnetic Monopole

A spin-1/2 degree of freedom $\mathbf{s}$ coupled via an exchange interaction J to a local magnetic field $\mathbf{S}$ is described by the two-level Hamiltonian

$$H(\mathbf{S}) = J\mathbf{S}\mathbf{s} = \frac{J}{2}\mathbf{S}\boldsymbol{\sigma} = \frac{J}{2} \begin{pmatrix} S_z & S_x - iS_y \\ S_x + iS_y & -S_z \end{pmatrix}, \quad (\text{A.176})$$

given in the basis $\{|\uparrow\rangle, |\downarrow\rangle\}$ of spin-up and spin-down eigenstates. For convenience we define $\mathbf{B} := J\mathbf{S}/2$ and write

$$H(\mathbf{B}) = \begin{pmatrix} B_z & B_x - iB_y \\ B_x + iB_y & -B_z \end{pmatrix}. \quad (\text{A.177})$$

The eigenenergies of Eq. (A.177) are then defined via

$$0 = \det \begin{pmatrix} B_z - E & B_x - iB_y \\ B_x + iB_y & -B_z - E \end{pmatrix} = E^2 - B_x^2 - B_y^2 - B_z^2, \quad (\text{A.178})$$

which yields

$$E_{\pm} = \pm \sqrt{B_x^2 + B_y^2 + B_z^2} =: \pm B. \quad (\text{A.179})$$

Note that these two eigenenergies become degenerate only for $B = 0$. The eigenvectors of $E_{\pm}$ are determined by

$$(H(\mathbf{B}) - E_{\pm})|\varphi_{\pm}(\mathbf{B})\rangle = (H(\mathbf{B}) - E_{\pm})(\varphi_{\pm,\uparrow}(\mathbf{B})|\uparrow\rangle + \varphi_{\pm,\downarrow}(\mathbf{B})|\downarrow\rangle) \stackrel{!}{=} 0. \quad (\text{A.180})$$

In the $\{|\uparrow\rangle, |\downarrow\rangle\}$ basis this becomes

$$\begin{pmatrix} B_z - E_{\pm} & B_x - iB_y \\ B_x + iB_y & -B_z - E_{\pm} \end{pmatrix} \begin{pmatrix} \varphi_{\pm,\uparrow}(\mathbf{B}) \\ \varphi_{\pm,\downarrow}(\mathbf{B}) \end{pmatrix} \stackrel{!}{=} \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \quad (\text{A.181})$$

which provides us with the condition

$$\varphi_{\pm,\downarrow}(\mathbf{B}) = \frac{(B_x + iB_y)}{(B_z + E_{\pm})} \varphi_{\pm,\uparrow}(\mathbf{B}) \quad (\text{A.182})$$

for the coefficients. We choose

$$|\varphi_{\pm}(\mathbf{B})\rangle = \frac{1}{(B_z + E_{\pm})} \begin{pmatrix} B_z + E_{\pm} \\ B_x + iB_y \end{pmatrix}. \quad (\text{A.183})$$

These eigenvectors are not yet normalised, so we compute their norm

$$\begin{aligned} |\varphi_{\pm}(\mathbf{B})| &= \sqrt{\frac{1}{(B_z + E_{\pm})^2} ((B_z + E_{\pm})^2 + (B_x - iB_y)(B_x + iB_y))} \\ &= \sqrt{\frac{1}{(B_z + E_{\pm})^2} (B_z^2 + 2B_zE_{\pm} + E_{\pm}^2 + B_x^2 + B_y^2)} \\ &= \sqrt{\frac{1}{(B_z + E_{\pm})^2} (2B_zE_{\pm} + E_{\pm}^2 + B^2)} \\ &= \sqrt{\frac{1}{(B_z + E_{\pm})^2} (2B_zE_{\pm} + E_{\pm}^2 + E_{\pm}^2)} \\ &= \sqrt{\frac{2E_{\pm}}{(B_z + E_{\pm})^2} (E_{\pm} + B_z)} \\ &= \sqrt{\frac{2E_{\pm}}{B_z + E_{\pm}}}, \end{aligned} \quad (\text{A.184})$$

and define the normalised eigenvectors

$$\begin{aligned}
|\psi_{\pm}(\mathbf{B})\rangle &= \frac{1}{|\varphi_{\pm}(\mathbf{B})|} |\varphi_{\pm}(\mathbf{B})\rangle \\
&= \sqrt{\frac{(B_z + E_{\pm})}{2E_{\pm}}} \frac{1}{(B_z + E_{\pm})} \begin{pmatrix} B_z + E_{\pm} \\ B_x + iB_y \end{pmatrix} \\
&= \frac{1}{\sqrt{2E_{\pm}(B_z + E_{\pm})}} \begin{pmatrix} B_z + E_{\pm} \\ B_x + iB_y \end{pmatrix}.
\end{aligned} \tag{A.185}$$

For $B \neq 0$, these states define two independent principal U(1) bundles $\psi_{\pm} \xrightarrow{\pi_{\pm}} (\mathbb{R}^3 \setminus \{\mathbf{0}\})$. We can use Eq. (A.185) to determine the Berry connections

$$\mathcal{A}_{\pm}(\mathbf{B}) = \langle \psi_{\pm}(\mathbf{B}) | d | \psi_{\pm}(\mathbf{B}) \rangle \tag{A.186}$$

of $\psi_{\pm} \xrightarrow{\pi_{\pm}} (\mathbb{R}^3 \setminus \{\mathbf{0}\})$ explicitly. In components, they read

$$\begin{aligned}
\mathcal{A}_{\pm}(\mathbf{B}) &= (\langle \uparrow | \psi_{\pm, \uparrow}(\mathbf{B})^* + \langle \downarrow | \psi_{\pm, \downarrow}(\mathbf{B})^*) \left[\sum_{\mu \in I} ((\partial_{\mu} \psi_{\pm, \uparrow}(\mathbf{B})) |\uparrow\rangle + (\partial_{\mu} \psi_{\pm, \downarrow}(\mathbf{B})) |\downarrow\rangle) dB_{\mu} \right] \\
&= \sum_{\mu \in I} (\psi_{\pm, \uparrow}(\mathbf{B})^* (\partial_{\mu} \psi_{\pm, \uparrow}(\mathbf{B})) + \psi_{\pm, \downarrow}(\mathbf{B})^* (\partial_{\mu} \psi_{\pm, \downarrow}(\mathbf{B}))) dB_{\mu},
\end{aligned} \tag{A.187}$$

where we introduced the index set $I = \{x, y, z\}$ to aid readability. We break this down in steps. To avoid clutter, we suppress explicit $\mathbf{B}$ -dependencies throughout the calculation, writing, for example, $E_{\pm}$ instead of $E_{\pm}(\mathbf{B})$. First, we compute

$$\partial_{\mu} E_{\pm} = \pm \partial_{\mu} \sqrt{B_x^2 + B_y^2 + B_z^2} = \pm \frac{2B_{\mu}}{2\sqrt{B_x^2 + B_y^2 + B_z^2}} = \frac{B_{\mu}}{E_{\pm}}, \tag{A.188}$$

where $\mu \in I$. Next, we find

$$\partial_{\mu} \sqrt{2E_{\pm}(B_z + E_{\pm})} = \frac{2\frac{B_{\mu}}{E_{\pm}}(B_z + E_{\pm}) + 2E_{\pm}\frac{B_{\mu}}{E_{\pm}}}{2\sqrt{2E_{\pm}(B_z + E_{\pm})}} = \frac{B_{\mu}(B_z + 2E_{\pm})}{E_{\pm}\sqrt{2E_{\pm}(B_z + E_{\pm})}} \tag{A.189}$$

for $\mu = x, y$ and

$$\partial_z \sqrt{2E_{\pm}(B_z + E_{\pm})} = \frac{2\frac{B_z}{E_{\pm}}(B_z + E_{\pm}) + 2E_{\pm}(1 + \frac{B_z}{E_{\pm}})}{2\sqrt{2E_{\pm}(B_z + E_{\pm})}} = \frac{(B_z + E_{\pm})^2}{E_{\pm}\sqrt{2E_{\pm}(B_z + E_{\pm})}} \tag{A.190}$$

for the partial derivative in z -direction. With these auxiliary results, we determine the derivatives of $\psi_{\pm, \uparrow}$ and $\psi_{\pm, \downarrow}$ separately. We begin with $\psi_{\pm, \uparrow}$ and get

$$\begin{aligned}
\partial_{\mu} \psi_{\pm, \uparrow} &= \partial_{\mu} \frac{(B_z + E_{\pm})}{\sqrt{2E_{\pm}(B_z + E_{\pm})}} \\
&= \partial_{\mu} \sqrt{\frac{(B_z + E_{\pm})}{2E_{\pm}}} \\
&= \frac{1}{2} \sqrt{\frac{2E_{\pm}}{(B_z + E_{\pm})}} \frac{\frac{B_{\mu}}{E_{\pm}} 2E_{\pm} - (B_z + E_{\pm}) 2\frac{B_{\mu}}{E_{\pm}}}{4E_{\pm}^2} \\
&= \sqrt{\frac{2E_{\pm}}{(B_z + E_{\pm})}} \frac{B_{\mu} - B_{\mu} - \frac{B_{\mu} B_z}{E_{\pm}}}{4E_{\pm}^2} \\
&= -\sqrt{\frac{2E_{\pm}}{(B_z + E_{\pm})}} \frac{B_{\mu} B_z}{4E_{\pm}^3} \\
&= -\frac{B_{\mu} B_z}{2E_{\pm}^2 \sqrt{2E_{\pm}(B_z + E_{\pm})}}
\end{aligned} \tag{A.191}$$

for $\mu = x, y$ and

$$\begin{aligned}
\partial_z \psi_{\pm, \uparrow} &= \partial_z \frac{(B_z + E_{\pm})}{\sqrt{2E_{\pm}(B_z + E_{\pm})}} \\
&= \partial_z \sqrt{\frac{(B_z + E_{\pm})}{2E_{\pm}}} \\
&= \frac{1}{2} \sqrt{\frac{2E_{\pm}}{(B_z + E_{\pm})}} \frac{(1 + \frac{B_{\mu}}{E_{\pm}})2E_{\pm} - (B_z + E_{\pm})2\frac{B_{\mu}}{E_{\pm}}}{4E_{\pm}^2} \\
&= \sqrt{\frac{2E_{\pm}}{(B_z + E_{\pm})}} \frac{B_z - B_z - \frac{B_z^2}{E_{\pm}} + E_{\pm}}{4E_{\pm}^2} \\
&= \sqrt{\frac{2E_{\pm}}{(B_z + E_{\pm})}} \frac{E_{\pm}^2 - B_z^2}{4E_{\pm}^3} \\
&= \frac{B_x^2 + B_y^2}{2E_{\pm}^2 \sqrt{2E_{\pm}(B_z + E_{\pm})}}
\end{aligned} \tag{A.192}$$

for the partial derivative in z -direction. Analogously, the partial derivatives of the $\psi_{\pm, \downarrow}$ component yield

$$\begin{aligned}
\partial_x \psi_{\pm, \downarrow} &= \partial_x \frac{(B_x + iB_y)}{\sqrt{2E_{\pm}(B_z + E_{\pm})}} \\
&= \frac{\sqrt{2E_{\pm}(B_z + E_{\pm})} - (B_x + iB_y) \frac{B_x(B_z + 2E_{\pm})}{E_{\pm} \sqrt{2E_{\pm}(B_z + E_{\pm})}}}{2E_{\pm}(B_z + E_{\pm})} \\
&= \frac{2E_{\pm}(B_z + E_{\pm})E_{\pm} - (B_x + iB_y)B_x(B_z + 2E_{\pm})}{E_{\pm}(2E_{\pm}(B_z + E_{\pm}))^{3/2}} \\
&= \frac{2(B_z^2 + B_y^2)(E_{\pm} + B_z) + B_x^2 B_z - iB_x B_y(B_z + 2E_{\pm})}{E_{\pm}(2E_{\pm}(B_z + E_{\pm}))^{3/2}}
\end{aligned} \tag{A.193}$$

in the x -direction,

$$\begin{aligned}
\partial_y \psi_{\pm, \downarrow} &= \partial_y \frac{(B_x + iB_y)}{\sqrt{2E_{\pm}(B_z + E_{\pm})}} \\
&= \frac{i\sqrt{2E_{\pm}(B_z + E_{\pm})} - (B_x + iB_y) \frac{B_y(B_z + 2E_{\pm})}{E_{\pm} \sqrt{2E_{\pm}(B_z + E_{\pm})}}}{2E_{\pm}(B_z + E_{\pm})} \\
&= \frac{i2E_{\pm}(B_z + E_{\pm})E_{\pm} - (B_x + iB_y)B_y(B_z + 2E_{\pm})}{E_{\pm}(2E_{\pm}(B_z + E_{\pm}))^{3/2}} \\
&= \frac{i2(B_z^2 + B_x^2)(E_{\pm} + B_z) + iB_y^2 B_z - B_x B_y(B_z + 2E_{\pm})}{E_{\pm}(2E_{\pm}(B_z + E_{\pm}))^{3/2}}
\end{aligned} \tag{A.194}$$

in the y -direction, and

$$\begin{aligned}
\partial_z \psi_{\pm, \downarrow} &= \partial_z \frac{(B_x + iB_y)}{\sqrt{2E_{\pm}(B_z + E_{\pm})}} \\
&= - \frac{(B_x + iB_y) \frac{(B_z + 2E_{\pm})^2}{E_{\pm} \sqrt{2E_{\pm}(B_z + E_{\pm})}}}{2E_{\pm}(B_z + E_{\pm})} \\
&= - \frac{(B_x + iB_y)(B_z + 2E_{\pm})^2}{E_{\pm}(2E_{\pm}(B_z + E_{\pm}))^{3/2}}
\end{aligned} \tag{A.195}$$

in the z -direction. If we plug these into Eq. (A.187) we get

$$\begin{aligned}
\mathcal{A}_\pm(\mathbf{B}) &= \sum_{\mu \in I} \left(\psi_{\pm, \uparrow}(\mathbf{B})^* (\partial_\mu \psi_{\pm, \uparrow}(\mathbf{B})) + \psi_{\pm, \downarrow}(\mathbf{B})^* (\partial_\mu \psi_{\pm, \downarrow}(\mathbf{B})) \right) d\mathbf{B}_\mu \\
&= \left(-\frac{(B_z + E_\pm)}{\sqrt{2E_\pm(B_z + E_\pm)}} \frac{B_x B_z}{2E_\pm^2 \sqrt{2E_\pm(B_z + E_\pm)}} \right. \\
&\quad \left. + \frac{(B_x - iB_y)}{\sqrt{2E_\pm(B_z + E_\pm)}} \frac{2(B_z^2 + B_y^2)(E_\pm + B_z) + B_x^2 B_z - iB_x B_y(B_z + 2E_\pm)}{E_\pm (2E_\pm(B_z + E_\pm))^{3/2}} \right) d\mathbf{B}_x \\
&\quad + \left(-\frac{(B_z + E_\pm)}{\sqrt{2E_\pm(B_z + E_\pm)}} \frac{B_y B_z}{2E_\pm^2 \sqrt{2E_\pm(B_z + E_\pm)}} \right. \\
&\quad \left. + \frac{(B_x - iB_y)}{\sqrt{2E_\pm(B_z + E_\pm)}} \frac{2i(B_z^2 + B_x^2)(E_\pm + B_z) + iB_y^2 B_z - B_x B_y(B_z + 2E_\pm)}{E_\pm (2E_\pm(B_z + E_\pm))^{3/2}} \right) d\mathbf{B}_y \\
&\quad + \left(\frac{(B_z + E_\pm)}{\sqrt{2E_\pm(B_z + E_\pm)}} \frac{B_x^2 + B_y^2}{2E_\pm^2 \sqrt{2E_\pm(B_z + E_\pm)}} \right. \\
&\quad \left. - \frac{(B_x - iB_y)}{\sqrt{2E_\pm(B_z + E_\pm)}} \frac{(B_x + iB_y)(B_z + E_\pm)^2}{E_\pm (2E_\pm(B_z + E_\pm))^{3/2}} \right) d\mathbf{B}_z \\
&= \left(\frac{-(B_z + E_\pm)^2 B_x B_z + (B_x - iB_y) (2(B_z^2 + B_y^2)(E_\pm + B_z) + B_x^2 B_z - iB_x B_y(B_z + 2E_\pm))}{4E_\pm^3 (B_z + E_\pm)^2} \right) d\mathbf{B}_x \\
&\quad + \left(\frac{-(B_z + E_\pm)^2 B_y B_z + (B_y + iB_x) (2(B_z^2 + B_x^2)(E_\pm + B_z) + B_y^2 B_z + iB_x B_y(B_z + 2E_\pm))}{4E_\pm^3 (B_z + E_\pm)^2} \right) d\mathbf{B}_y \\
&\quad + \left(\frac{(B_z + E_\pm)^2 (B_x^2 + B_y^2) - (B_z + E_\pm)^2 (B_x^2 + B_y^2)}{4E_\pm^3 (B_z + E_\pm)^2} \right) d\mathbf{B}_z \\
&=: \mathcal{A}_{\pm, x}(\mathbf{B}) d\mathbf{B}_x + \mathcal{A}_{\pm, y}(\mathbf{B}) d\mathbf{B}_y, \tag{A.196}
\end{aligned}$$

where we defined the coefficients

$$\begin{aligned}
\mathcal{A}_{\pm, x}(\mathbf{B}) &:= \frac{-(B_z + E_\pm)^2 B_x B_z + (B_x - iB_y) (2(B_z^2 + B_y^2)(E_\pm + B_z) + B_x^2 B_z - iB_x B_y(B_z + 2E_\pm))}{4E_\pm^3 (B_z + E_\pm)^2} \\
\mathcal{A}_{\pm, y}(\mathbf{B}) &:= \frac{-(B_z + E_\pm)^2 B_y B_z + (B_y + iB_x) (2(B_z^2 + B_x^2)(E_\pm + B_z) + B_y^2 B_z + iB_x B_y(B_z + 2E_\pm))}{4E_\pm^3 (B_z + E_\pm)^2} \tag{A.197}
\end{aligned}$$

of the Berry connection. Note that the $\mathcal{A}_{\pm, x}(\mathbf{B})$ and $\mathcal{A}_{\pm, y}(\mathbf{B})$ are the same up to a coordinate rotation $B_x \mapsto B_y$ and $B_y \mapsto -B_x$. This allows us to compute $\mathcal{A}_{\pm, x}(\mathbf{B})$ and $\mathcal{A}_{\pm, y}(\mathbf{B})$ simultaneously. Specifically, we compute $\mathcal{A}_{\pm, x}(\mathbf{B})$ first,

$$\begin{aligned}
\mathcal{A}_{\pm, x}(\mathbf{B}) &= \frac{-(B_z + E_\pm)^2 B_x B_z + (B_x - iB_y) (2(B_z^2 + B_y^2)(E_\pm + B_z) + B_x^2 B_z - iB_x B_y(B_z + 2E_\pm))}{4E_\pm^3 (B_z + E_\pm)^2} \\
&= \left(\frac{-2i(E_\pm + B_z)B_y^3 + (-\cancel{B_x B_z^3} + \cancel{2B_x E_\pm} + \cancel{2B_x B_z} - \cancel{B_x B_z} - \cancel{2B_z E_\pm})B_y^2}{4E_\pm^3 (B_z + E_\pm)^2} \right. \\
&\quad + \frac{(-2iB_x^2 B_z - 2iB_x^2 E_\pm - 2iB_z^2 E_\pm - 2iB_z^2 B_z)B_y}{4E_\pm^3 (B_z + E_\pm)^2} \\
&\quad \left. + \frac{(-\cancel{B_z^3 B_x} - \cancel{2B_z^2 B_x E_\pm} - \cancel{B_z^3 B_x} - \cancel{B_x B_z^3} + \cancel{2B_x B_z^2 E_\pm} + \cancel{2B_x B_z^3} + \cancel{B_z^3 B_z})}{4E_\pm^3 (B_z + E_\pm)^2} \right) \\
&= -\frac{2i(E_\pm + B_z) (B_y^3 + (B_x^2 + B_z^2)B_y)}{4E_\pm^3 (B_z + E_\pm)^2} \\
&= -\frac{i((B_y^2 + B_x^2 + B_z^2)B_y)}{2E_\pm^3 (B_z + E_\pm)} \\
&\stackrel{(\circ)}{=} -\frac{iE_\pm^2 B_y}{2E_\pm^3 (B_z + E_\pm)} \\
&= -\frac{iB_y}{2E_\pm (B_z + E_\pm)}, \tag{A.198}
\end{aligned}$$

where we have used $E_{\pm}^2 = B^2 = B_x^2 + B_y^2 + B_z^2$ in $(\diamond)$. Then we perform the coordinate rotation $B_x \mapsto B_y$ and $B_y \mapsto -B_x$ on the result to get

$$\mathcal{A}_{\pm,y}(\mathbf{B}) = \mp \frac{iB_x}{2E_{\pm}(B_z + E_{\pm})}, \quad (\text{A.199})$$

where we used that $E_{\pm}$ only features B_y^2 so that $E_{\pm} \mapsto E_{\pm}$ under the coordinate transformation. Combined, we find

$$\begin{aligned} \mathcal{A}_{\pm}(\mathbf{B}) &= \mathcal{A}_{\pm,x}(\mathbf{B})dB_x + \mathcal{A}_{\pm,y}(\mathbf{B})dB_y \\ &= -\frac{i}{2} \frac{B_y dB_x - B_x dB_y}{E_{\pm}(B_z + E_{\pm})}. \end{aligned} \quad (\text{A.200})$$

If we plug in Eq. (A.179) we get

$$\mathcal{A}_{\pm}(\mathbf{B}) = \mp \frac{i}{2} \frac{B_y dB_x - B_x dB_y}{B(B_z \pm B)}. \quad (\text{A.201})$$

This tells us that in the current gauge choice, i.e. in the current choice Eq. (A.185) of eigenstates, the Berry connections $\mathcal{A}_{\pm}(\mathbf{B})$ of the $\psi_{\pm} \xrightarrow{\pi_{\pm}} (\mathbb{R}^3 \setminus \{\mathbf{0}\})$ bundles are only well-defined if $B \neq \mp B_z$. Specifically, $\mathcal{A}_{+}(\mathbf{B})$ is only well-defined away from the negative z -axis, while $\mathcal{A}_{-}(\mathbf{B})$ is only well-defined away from the positive z -axis. This singularity in the Berry connection is famously known as the *Dirac string* [2]. The location of the Dirac string can be manipulated through gauge transformations. However, it can never be removed. This is a prototypical example of a topological obstruction: as we will see shortly, the first Chern number of the principal U(1) bundles $\psi_{\pm} \xrightarrow{\pi_{\pm}} (\mathbb{R}^3 \setminus \{\mathbf{0}\})$ is non-zero. We mentioned in Sec. 2.3 that a non-zero first Chern number obstructs the definition of a nowhere-vanishing global section. In the associated principal U(1) bundle this manifests as a singularity because the U(1) phase of the zero vector (which necessarily exists somewhere as no nowhere-vanishing global sections exist) is ill-defined. To see that the Chern numbers of the $\psi_{\pm} \xrightarrow{\pi_{\pm}} (\mathbb{R}^3 \setminus \{\mathbf{0}\})$ bundles are non-zero, we compute the Berry curvature

$$\begin{aligned} \mathcal{F}_{\pm}(\mathbf{B}) &= d\mathcal{A}_{\pm}(\mathbf{B}) \\ &= \sum_{\mu,\nu \in I} \partial_{\mu} \mathcal{A}_{\pm,\nu}(\mathbf{B}) dB_{\mu} \wedge dB_{\nu} \\ &= \partial_x \mathcal{A}_{\pm,y}(\mathbf{B}) dB_x \wedge dB_y + \partial_y \mathcal{A}_{\pm,x}(\mathbf{B}) dB_y \wedge dB_x \\ &\quad + \partial_z \mathcal{A}_{\pm,x}(\mathbf{B}) dB_z \wedge dB_x + \partial_z \mathcal{A}_{\pm,y}(\mathbf{B}) dB_z \wedge dB_y \\ &= \partial_x \left(\frac{iB_x}{2E_{\pm}(B_z + E_{\pm})} \right) dB_x \wedge dB_y + \partial_y \left(-\frac{iB_y}{2E_{\pm}(B_z + E_{\pm})} \right) dB_y \wedge dB_x \\ &\quad + \partial_z \left(-\frac{iB_y}{2E_{\pm}(B_z + E_{\pm})} \right) dB_z \wedge dB_x + \partial_z \left(\frac{iB_x}{2E_{\pm}(B_z + E_{\pm})} \right) dB_z \wedge dB_y \\ &\stackrel{(\diamond)}{=} \left[\partial_x \left(\frac{iB_x}{2E_{\pm}(B_z + E_{\pm})} \right) + \partial_y \left(\frac{iB_y}{2E_{\pm}(B_z + E_{\pm})} \right) \right] dB_x \wedge dB_y \\ &\quad + \partial_z \left(\frac{-iB_y}{2E_{\pm}(B_z + E_{\pm})} \right) dB_z \wedge dB_x - \partial_z \left(\frac{iB_x}{2E_{\pm}(B_z + E_{\pm})} \right) dB_y \wedge dB_z \\ &=: \mathcal{F}_{\pm,xy}(\mathbf{B}) dB_x \wedge dB_y + \mathcal{F}_{\pm,zx}(\mathbf{B}) dB_z \wedge dB_x + \mathcal{F}_{\pm,yz}(\mathbf{B}) dB_y \wedge dB_z, \end{aligned} \quad (\text{A.202})$$

where we defined the Berry curvature coefficients $\mathcal{F}_{\pm,\mu\nu}(\mathbf{B})$ with respect to the standard order xyz of indices. In $(\diamond)$ we used $dB_y \wedge dB_x = -dB_x \wedge dB_y$ to combine the first two terms into $\mathcal{F}_{\pm,xy}(\mathbf{B})$. In order to compute this explicitly, we first use Eq. (A.188) to compute the auxiliary result

$$\begin{aligned} \partial_{\mu} \frac{B_{\nu}}{2(\pm B)(B_z \pm B)} &= \pm \frac{\delta_{\mu\nu} 2B(B_z \pm B) - B_{\nu} \left(2\frac{B_{\mu}}{B}(B_z \pm B) + 2B(\delta_{\mu z} \pm \frac{B_{\mu}}{B}) \right)}{4B^2(B_z \pm B)^2} \\ &= \pm \frac{\delta_{\mu\nu} 2B^2(B_z \pm B) - B_{\nu} (2B_{\mu}(B_z \pm B) + 2B(\delta_{\mu z} B \pm B_{\mu}))}{4B^3(B_z \pm B)^2} \\ &= \pm \frac{\delta_{\mu\nu} 2B^2(B_z \pm B) - B_{\nu} (2B_{\mu}(B_z \pm 2B) + 2\delta_{\mu z} B^2)}{4B^3(B_z \pm B)^2}. \end{aligned} \quad (\text{A.203})$$

With this we get

$$\begin{aligned}
\partial_x \mathcal{A}_{\pm,y}(\mathbf{B}) &= \partial_x \left(\frac{iB_x}{2E_{\pm}(B_z + E_{\pm})} \right) \\
&= \pm i \left[\frac{2B^2(B_z \pm B) - 2B_x^2(B_z \pm 2B)}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{(2B^2 - 2B_x^2)(B_z \pm B) \mp 2B_x^2 B}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{(2B_y^2 + 2B_z^2)(B_z \pm B) \mp 2B_x^2 B}{4B^3(B_z \pm B)^2} \right], \tag{A.204}
\end{aligned}$$

and

$$\begin{aligned}
\partial_y \mathcal{A}_{\pm,x}(\mathbf{B}) &= \partial_y \left(\frac{-iB_y}{2E_{\pm}(B_z + E_{\pm})} \right) \\
&= \mp i \left[\frac{2B^2(B_z \pm B) - 2B_y^2(B_z \pm 2B)}{4B^3(B_z \pm B)^2} \right] \\
&= \mp i \left[\frac{(2B^2 - 2B_y^2)(B_z \pm B) \mp 2B_y^2 B}{4B^3(B_z \pm B)^2} \right] \\
&= \mp i \left[\frac{(2B_x^2 + 2B_z^2)(B_z \pm B) \mp 2B_y^2 B}{4B^3(B_z \pm B)^2} \right], \tag{A.205}
\end{aligned}$$

for the partial derivatives in the x - and y -directions. The partial derivatives in the z -direction immediately yield the corresponding coefficients of the Berry curvature – they become

$$\begin{aligned}
\mathcal{F}_{\pm,yz}(\mathbf{B}) &= -\partial_z \mathcal{A}_{\pm,y}(\mathbf{B}) = -\partial_z \left(\frac{iB_x}{2E_{\pm}(B_z + E_{\pm})} \right) \\
&= \pm i \left[\frac{B_x(2B_z(B_z \pm 2B) + 2B^2)}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{B_x(2B_z^2 \pm 4B_z B + 2B^2)}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{2B_x(B_z \pm B)^2}{4B^3(B_z \pm B)^2} \right] \\
&= \pm \frac{iB_x}{2B^3}, \tag{A.206}
\end{aligned}$$

and

$$\begin{aligned}
\mathcal{F}_{\pm,zx}(\mathbf{B}) &= \partial_z \mathcal{A}_{\pm,x}(\mathbf{B}) = \partial_z \left(\frac{-iB_y}{2E_{\pm}(B_z + E_{\pm})} \right) \\
&= \pm i \left[\frac{B_y(2B_z(B_z \pm 2B) + 2B^2)}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{B_y(2B_z^2 \pm 4B_z B + 2B^2)}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{2B_y(B_z \pm B)^2}{4B^3(B_z \pm B)^2} \right] \\
&= \pm \frac{iB_y}{2B^3}, \tag{A.207}
\end{aligned}$$

respectively. Next, we determine the combined coefficient $\mathcal{F}_{\pm,xy}(\mathbf{B})$ of $dB_x \wedge dB_y$ in Eq. (A.202).

We find

$$\begin{aligned}
\mathcal{F}_{\pm,xy}(\mathbf{B}) &= \partial_x \left(\frac{iB_x}{2E_{\pm}(B_z + E_{\pm})} \right) + \partial_y \left(\frac{iB_y}{2E_{\pm}(B_z + E_{\pm})} \right) \\
&= \pm i \left[\frac{(2B_y^2 + 2B_z^2)(B_z \pm B) \mp 2B_x^2 B}{4B^3(B_z \pm B)^2} + \frac{(2B_x^2 + 2B_z^2)(B_z \pm B) \mp 2B_y^2 B}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{(2B_y^2 + 2B_z^2)(B_z \pm B) \mp 2B_x^2 B + (2B_x^2 + 2B_z^2)(B_z \pm B) \mp 2B_y^2 B}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{(2B_x^2 + 2B_y^2 + 2B_z^2)(B_z \pm B) \mp 2B_x^2 B + 2B_z^2(B_z \pm B) \mp 2B_y^2 B}{4B^3(B_z \pm B)^2} \right] \\
&\stackrel{(\diamond)}{=} \pm i \left[\frac{2B^2(B_z \pm B) \mp (2B_x^2 + 2B_y^2 + 2B_z^2)B + 2B_z^2(B_z \pm 2B)}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{2B^2(B_z \pm B) \mp 2B^2 \mp 2B_z^2(B_z \pm 2B)}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{2B^2 B_z + 2B_z^2 B_z \pm 4B_z^2 B}{4B^3(B_z \pm 2B)^2} \right] \\
&= \pm i \left[\frac{2B_z(B^2 + B_z^2 \pm 2B_z B)}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \left[\frac{B_z 2(B_z \pm B)^2}{4B^3(B_z \pm B)^2} \right] \\
&= \pm i \frac{B_z}{2B^3}.
\end{aligned} \tag{A.208}$$

If we plug Eqs. (A.208), (A.207), and (A.206) into Eq. (A.202) we arrive at the simple expression

$$\mathcal{F}_{\pm}(\mathbf{B}) = \pm \frac{i}{2} \frac{B_z dB_x \wedge dB_y + B_y dB_z \wedge dB_x + B_x dB_y \wedge dB_z}{B^3} \tag{A.209}$$

for the Berry curvature of the magnetic monopole, cf. e.g. Ref. [39]. Several remarks are in order.

The first one is that one can rewrite the Berry connection $\mathcal{A}_{\pm}(\mathbf{B})$ from Eq. (A.201) as

$$\begin{aligned}
\mathcal{A}_{\pm}(\mathbf{B}) &= \mp \frac{i}{2} \frac{B_y dB_x - B_x dB_y}{B(B_z \pm B)} \\
&= \mp \frac{i}{2B^2} \frac{\begin{pmatrix} B_y & -B_x & 0 \end{pmatrix}}{\begin{pmatrix} B_z \\ B \end{pmatrix} \pm 1} \begin{pmatrix} dB_x \\ dB_y \\ dB_z \end{pmatrix} \\
&= \mp \frac{i}{2B^2} \frac{[\mathbf{B} \times \mathbf{e}_z]^{\top}}{(\mathbf{e}_z \mathbf{B} / B \pm 1)} d\mathbf{B} \\
&= \pm \frac{i}{2B^2} \frac{[\mathbf{e}_z \times \mathbf{B}]^{\top}}{(\mathbf{e}_z \mathbf{B} / B \pm 1)} d\mathbf{B} \\
&=: \mathbf{A}_{\pm}(\mathbf{B})^{\top} d\mathbf{B},
\end{aligned} \tag{A.210}$$

where we identified the numerator as an inner product between the two vectors $[\mathbf{B} \times \mathbf{e}_z]^{\top} = (B_y \quad -B_x \quad 0)$ and $d\mathbf{B} \equiv (dB_x \quad dB_y \quad dB_z)^{\top}$ and defined the coefficient field

$$\mathbf{A}_{\pm}(\mathbf{B}) = \pm \frac{i}{2B^2} \frac{\mathbf{e}_z \times \mathbf{B}}{(\mathbf{e}_z \mathbf{B} / B \pm 1)}. \tag{A.211}$$

This form of the Berry connection $\mathcal{A}_{\pm}(\mathbf{B})$ is more familiar to most physicists and makes the analogy with the electromagnetic gauge potential more transparent. Note that it is usually the coefficient field $\mathbf{A}_{\pm}(\mathbf{B})$ that is quoted as the electromagnetic gauge potential.

Along the same lines one can rewrite the Berry curvature $\mathcal{F}_\pm(\mathbf{B})$ from Eq. (A.209) as

$$\begin{aligned}
\mathcal{F}_\pm(\mathbf{B}) &= \pm \frac{i}{2} \frac{B_z dB_x \wedge dB_y + B_y dB_z \wedge dB_x + B_x dB_y \wedge dB_z}{B^3} \\
&= \pm \frac{i}{B^3} \begin{pmatrix} B_x & B_y & B_z \end{pmatrix} \begin{pmatrix} dB_y \wedge dB_z \\ dB_z \wedge dB_x \\ dB_x \wedge dB_y \end{pmatrix} \\
&= \pm \frac{i \mathbf{B}^\top}{B^3} [\star d\mathbf{B}] \\
&=: \mathbf{F}_\pm(\mathbf{B})^\top d\mathbf{S},
\end{aligned} \tag{A.212}$$

where we used the hodge operator $\star : \bigwedge^p T^*M \rightarrow \bigwedge^{n-p} T^*M$ to translate the vector $d\mathbf{B}$ of one-forms into a vector

$$d\mathbf{S} := \star d\mathbf{B} \equiv \begin{pmatrix} \star dB_x \\ \star dB_y \\ \star dB_z \end{pmatrix} = \begin{pmatrix} dB_y \wedge dB_z \\ dB_z \wedge dB_x \\ dB_x \wedge dB_y \end{pmatrix} \tag{A.213}$$

of two-forms. We also defined a coefficient field

$$\mathbf{F}_\pm(\mathbf{B}) = \pm \frac{i \mathbf{B}}{B^3} \tag{A.214}$$

of the Berry curvature. By definition, this is equal to the curl

$$\mathbf{F}_\pm(\mathbf{B}) = \text{curl} \mathbf{A}_\pm(\mathbf{B}) = \nabla \times \mathbf{A}_\pm(\mathbf{B}) \tag{A.215}$$

of the coefficient field $\mathbf{A}_\pm(\mathbf{B})$ associated with the Berry connection. This can be seen through explicit calculation (see above) or via the fundamental definition of the exterior derivative: if we write

$$\mathcal{A}_\pm(\mathbf{B}) = \mathbf{A}_\pm(\mathbf{B})^\top d\mathbf{B} = A_{\pm,x}(\mathbf{B}) dB_x + A_{\pm,y}(\mathbf{B}) dB_y + A_{\pm,z}(\mathbf{B}) dB_z, \tag{A.216}$$

then we automatically get

$$\begin{aligned}
\mathcal{F}_\pm(\mathbf{B}) &= d\mathcal{A}_\pm(\mathbf{B}) \\
&= (\partial_x A_{\pm,y}(\mathbf{B}) - \partial_y A_{\pm,x}(\mathbf{B})) dB_x \wedge dB_y \\
&\quad + (\partial_z A_{\pm,x}(\mathbf{B}) - \partial_x A_{\pm,z}(\mathbf{B})) dB_z \wedge dB_x \\
&\quad + (\partial_y A_{\pm,z}(\mathbf{B}) - \partial_z A_{\pm,y}(\mathbf{B})) dB_y \wedge dB_z \\
&= \mathbf{F}_\pm(\mathbf{B})^\top d\mathbf{S},
\end{aligned} \tag{A.217}$$

and hence the coefficient vector

$$\mathbf{F}_\pm(\mathbf{B}) = \begin{pmatrix} \partial_x A_{\pm,y}(\mathbf{B}) - \partial_y A_{\pm,x}(\mathbf{B}) \\ \partial_z A_{\pm,x}(\mathbf{B}) - \partial_x A_{\pm,z}(\mathbf{B}) \\ \partial_y A_{\pm,z}(\mathbf{B}) - \partial_z A_{\pm,y}(\mathbf{B}) \end{pmatrix} = \text{curl} \mathbf{A}_\pm(\mathbf{B}), \tag{A.218}$$

by the skew-symmetry of the wedge product. The curvature $\mathcal{F}_\pm(\mathbf{B})$, and in the physics literature often its coefficient field $\mathbf{F}_\pm(\mathbf{B})$, is then usually compared to the magnetic field. The form of the gauge potential in Eq. (A.211) and the curvature vector field in Eq. (A.214) are strongly reminiscent of the gauge potential and magnetic field of a magnetic Dirac monopole.

It is worth conceding that the vector of two-forms in Eq. (A.213) is an unusual sight. In physics, we are accustomed to expressions like $\int_S \mathbf{F} d\mathbf{S}$ representing the flux of a magnetic field $\mathbf{F}$ through a surface S . This leads to shorthand statements like “integrating a magnetic field over a surface.” However, the mathematical theory of integration on manifolds does not assign meaning to the integration of vector fields over surfaces. The proper objects for integration over two-dimensional manifolds are two-forms. To translate these kinds of shorthand statements into the language of differential forms, we typically rely on two central tools: the musical isomorphism $\flat : TM \rightarrow T^*M$, which maps vector fields $X = X_\mu \partial_\mu$ on an n -dimensional manifold M to one-forms $x = X_\mu dx^\mu$, and the Hodge star, which maps p -forms on M to

$(n - p)$ -forms. If we choose to understand the curvature coefficient field $\mathbf{F}_\pm(\mathbf{B})$ from Eq. (A.214) as a vector field

$$\mathbf{F}_\pm(\mathbf{B}) = \mathbf{F}_{\pm,x}(\mathbf{B})\partial_x + \mathbf{F}_{\pm,y}(\mathbf{B})\partial_y + \mathbf{F}_{\pm,z}(\mathbf{B})\partial_z = \frac{iB_x}{B^3}\partial_x + \frac{iB_y}{B^3}\partial_y + \frac{iB_z}{B^3}\partial_z, \quad (\text{A.219})$$

then its integral over a two-dimensional manifold $S \in \mathbb{R}^3 \setminus \{\mathbf{0}\}$ corresponds to the expression

$$\begin{aligned} \int_S \star \mathbf{F}_\pm(\mathbf{B})^\flat &= \int_S \star \left(\frac{iB_x}{B^3}\partial_x + \frac{iB_y}{B^3}\partial_y + \frac{iB_z}{B^3}\partial_z \right)^\flat \\ &= \int_S \star \left(\frac{iB_x}{B^3}dB_x + \frac{iB_y}{B^3}dB_y + \frac{iB_z}{B^3}dB_z \right) \\ &= \int_S \left(\frac{iB_x}{B^3}dB_y \wedge dB_z + \frac{iB_y}{B^3}dB_z \wedge dB_x + \frac{iB_z}{B^3}dB_x \wedge dB_y \right) \\ &= \int_S \mathbf{F}_\pm(\mathbf{B})^\top d\mathbf{S}, \end{aligned} \quad (\text{A.220})$$

and hence the integral of the Berry curvature two-form from Eq. (A.212).

The second remark is that one can integrate the Berry curvature $\mathcal{F}_\pm(\mathbf{B})$ over any closed two-dimensional surface $S \subset \mathbb{R}^3 \setminus \{\mathbf{0}\}$ to get a Chern characterising the principal $U(1)$ bundle over S . Choose, for instance, the standard embedding $\{\mathbf{B} \in \mathbb{R}^3 \setminus \{\mathbf{0}\} \mid |\mathbf{B}| = 1\}$ of two-sphere $\mathbb{S}^2 \subset \mathbb{R}^3 \setminus \{\mathbf{0}\}$. To compute the Chern number over $\mathbb{S}^2$ we assume spherical coordinates

$$\begin{aligned} B_x &= B \sin \theta \cos \phi \\ B_y &= B \sin \theta \sin \phi \\ B_z &= B \cos \theta, \end{aligned} \quad (\text{A.221})$$

where $\phi \in [0, 2\pi]$ and $\theta \in [0, \pi]$. With this, we find

$$\begin{aligned} dB_x &= \sin \theta \cos \phi dB - B \cos \theta \cos \phi d\theta - B \sin \theta \sin \phi d\phi \\ dB_y &= \sin \theta \sin \phi dB + B \cos \theta \sin \phi d\theta + B \sin \theta \cos \phi d\phi, \end{aligned} \quad (\text{A.222})$$

such that the Berry connection becomes

$$\begin{aligned} \mathcal{A}_\pm(\mathbf{B}) &= \mp \frac{i}{2} \frac{B_y dB_x - B_x dB_y}{B(B_z \pm B)} \\ &= \frac{i}{2} \left[\frac{(B \sin \theta \sin \phi)([\sin \theta \cos \phi dB + B \cos \theta \cos \phi d\theta] - B \sin \theta \sin \phi d\phi)}{(\pm B)(\pm B + B \cos \theta)} \right. \\ &\quad \left. - \frac{(B \sin \theta \cos \phi)([\sin \theta \sin \phi dB + B \cos \theta \sin \phi d\theta] + B \sin \theta \cos \phi d\phi)}{B^2(1 \pm B \cos \theta)} \right] \\ &= \frac{i}{2} \frac{B^2(\sin^2 \theta \sin^2 \phi + \sin^2 \theta \cos^2 \phi)}{B^2(1 \pm \cos \theta)} d\phi \\ &= \frac{i}{2} \frac{\sin^2 \theta}{(1 \pm \cos \theta)} d\phi \\ &= \frac{i}{2} \frac{(1 - \cos^2 \theta)}{(1 \pm \cos \theta)} d\phi \\ &= \frac{i}{2} \frac{(1 + \cos \theta)(1 - \cos \theta)}{(1 \pm \cos \theta)} d\phi \\ &= \frac{i}{2} (1 \mp \cos \theta) d\phi. \end{aligned} \quad (\text{A.223})$$

If we compute the Berry curvature from this we obtain a similarly simple expression

$$\mathcal{F}_\pm(\mathbf{B}) = d \left[\frac{i}{2} (1 \mp \cos \theta) d\phi \right] = \pm \frac{i}{2} \sin \theta d\theta \wedge d\phi. \quad (\text{A.224})$$

The first Chern number of the principal $U(1)$ bundle over $\mathbb{S}^2 \subset \mathbb{R}^3 \setminus \{\mathbf{0}\}$ is then

$$\begin{aligned}
C_1(\mathcal{F}_\pm(\mathbf{B})) &= \langle [c_1(\mathcal{F}_\pm(\mathbf{B}))], [\mathbb{S}^2] \rangle = \frac{i}{2\pi} \int_{\mathbb{S}^2} \mathcal{F}_\pm(\mathbf{B}) \\
&= \frac{i}{2\pi} \int_0^\pi \int_0^{2\pi} \left[\pm \frac{i}{2} \sin \theta \, d\theta \wedge d\phi \right] \\
&= \mp \frac{1}{4\pi} \int_0^\pi \sin \theta \, d\theta \, d\phi \\
&= \mp \frac{2\pi}{4\pi} \int_0^\pi \sin \theta \, d\theta \\
&= \mp \frac{1}{2} \left[-\cos \theta \right]_0^\pi \\
&= \mp 1.
\end{aligned} \tag{A.225}$$

Note that the first Chern number is also equal to the first Chern character, i.e.

$$C_1(\mathcal{F}_\pm(\mathbf{B})) = Ch_1(\mathcal{F}_\pm(\mathbf{B})). \tag{A.226}$$

The fact that this Chern number is non-trivial tells us that there exists no nowhere-vanishing global section. There always has to be at least one point where the section is zero and the phase is ill-defined. Intuitively speaking, this point is where the Dirac string of the monopole that sits at $\mathbf{0} \in \mathbb{R}^3$ punctures the enveloping two-dimensional surface $S \subset \mathbb{R}^3 \setminus \{\mathbf{0}\}$, here: $S = \mathbb{S}^2$.

A.10 TRS, Fourier Transform and Topology of the Kane–Mele Model

Here, we provide a collection of explicit calculations for Chap. 8.

TRS Invariance of the Kane–Mele Hamiltonian

First we show that the Kane–Mele Hamiltonian from Eq. (8.1) is invariant under the TRS transformation Eq. (3.15) of spin one-half fermions:

$$\begin{aligned}
\mathcal{T}H_{\text{KM}}\mathcal{T}^\dagger &= \mathcal{T} \left(-t_{\text{hop}} \sum_{\langle j,k \rangle_\alpha} c_{j\alpha}^\dagger c_{k\alpha} + V \sum_{j,\alpha} \varepsilon_j c_{j\alpha}^\dagger c_{j\alpha} + it_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\beta}} \nu_{jk} \sigma_z^{\alpha\beta} c_{j\alpha}^\dagger c_{k\beta} \right) \mathcal{T}^\dagger \\
&= -t_{\text{hop}} \sum_{\langle j,k \rangle_\alpha} \mathcal{T} c_{j\alpha}^\dagger \mathcal{T} \mathcal{T}^\dagger c_{k\alpha} \mathcal{T}^\dagger + V \sum_{j,\alpha} \varepsilon_j \mathcal{T} c_{j\alpha}^\dagger \mathcal{T}^\dagger \mathcal{T} c_{j\alpha} \mathcal{T}^\dagger - it_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\beta}} \nu_{jk} \sigma_z^{\alpha\beta} \mathcal{T} c_{j\alpha}^\dagger \mathcal{T}^\dagger \mathcal{T} c_{k\beta} \mathcal{T}^\dagger \\
&\stackrel{(\diamond)}{=} t_{\text{hop}} \sum_{\langle j,k \rangle_{\alpha,\beta,\gamma}} \sigma_y^{\alpha\beta} \sigma_y^{\alpha\gamma} c_{j\beta}^\dagger c_{k\gamma} - V \sum_{j_{\alpha,\beta,\gamma}} \varepsilon_j \sigma_y^{\alpha\beta} \sigma_y^{\alpha\gamma} c_{j\beta}^\dagger c_{j\gamma} + it_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\beta,\gamma,\eta}} \nu_{jk} \sigma_z^{\alpha\beta} \sigma_y^{\alpha\gamma} \sigma_y^{\beta\eta} c_{j\gamma}^\dagger c_{k\eta} \\
&\stackrel{(*)}{=} -t_{\text{hop}} \sum_{\langle j,k \rangle_{\alpha,\beta,\gamma}} \sigma_y^{\beta\alpha} \sigma_y^{\alpha\gamma} c_{j\beta}^\dagger c_{k\gamma} + V \sum_{j_{\alpha,\beta,\gamma}} \varepsilon_j \sigma_y^{\beta\alpha} \sigma_y^{\alpha\gamma} c_{j\beta}^\dagger c_{j\gamma} - it_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\beta,\gamma,\eta}} \nu_{jk} \sigma_y^{\gamma\alpha} (\sigma_z^{\alpha\beta} \sigma_y^{\beta\eta}) c_{j\gamma}^\dagger c_{k\eta} \\
&\stackrel{(*)}{=} -t_{\text{hop}} \sum_{\langle j,k \rangle_{\beta,\gamma}} \delta_{\beta\gamma} c_{j\beta}^\dagger c_{k\gamma} + V \sum_{j_{\beta,\gamma}} \varepsilon_j \delta_{\beta\gamma} c_{j\beta}^\dagger c_{j\gamma} - it_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\gamma,\eta}} \nu_{jk} \sigma_y^{\gamma\alpha} (-i\sigma_x^{\alpha\eta}) c_{j\gamma}^\dagger c_{k\eta} \\
&= -t_{\text{hop}} \sum_{\langle j,k \rangle_\beta} c_{j\beta}^\dagger c_{k\beta} + V \sum_{j,\beta} \varepsilon_j c_{j\beta}^\dagger c_{j\beta} - t_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\gamma,\eta}} \nu_{jk} (\sigma_y^{\gamma\alpha} \sigma_x^{\alpha\eta}) c_{j\gamma}^\dagger c_{k\eta} \\
&= -t_{\text{hop}} \sum_{\langle j,k \rangle_\alpha} c_{j\alpha}^\dagger c_{k\alpha} + V \sum_{j,\alpha} \varepsilon_j c_{j\alpha}^\dagger c_{j\alpha} - t_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle_{\gamma,\eta}} \nu_{jk} (-i\sigma_z^{\gamma\eta}) c_{j\gamma}^\dagger c_{k\eta} \\
&= -t_{\text{hop}} \sum_{\langle j,k \rangle_\alpha} c_{j\alpha}^\dagger c_{k\alpha} + V \sum_{j,\alpha} \varepsilon_j c_{j\alpha}^\dagger c_{j\alpha} + it_{\text{SO}} \sum_{\langle\langle j,k \rangle\rangle_{\alpha,\beta}} \nu_{jk} \sigma_z^{\alpha\beta} c_{j\alpha}^\dagger c_{k\beta} \\
&= H_{\text{KM}}.
\end{aligned} \tag{A.227}$$

Here, we plugged in Eq. (3.15) in $(\diamond)$, applied the skew-symmetry $\sigma_y^\dagger = -\sigma_y$ of the y -Pauli matrix in $(*)$, and made repeated use of the relation $\sigma_j \sigma_k = \delta_{jk} + i\epsilon_{jkl} \sigma_l$ for the product of Pauli matrices from $(*)$ onwards.

Diagonalisation of the Kane–Mele Model in k -Space

We explicitly diagonalise Eq. (8.1) using the Fourier transform $c_{j\alpha} = 1/\sqrt{L} \sum_{\mathbf{k}} e^{i\mathbf{k}\mathbf{R}} c_{\mathbf{k}\alpha}$ of the elementary annihilation and creation operators. Like in the Haldane model, the honeycomb lattice structure causes some subtleties in combination with the SOC hopping $\propto it_{\text{SO}}$ of the Kane–Mele model. To deal with this, we formally separate the elementary field operators into sublattice- A operators a_j and sublattice- B operators b_j , writing

$$\begin{aligned}
H_{\text{KM}} &= -t_{\text{hop}} \sum_{j,\alpha} \sum_{k=1}^3 [a_{j\alpha}^\dagger b_{j+\nu_k\alpha} + b_{j+\nu_k\alpha}^\dagger a_{j\alpha}] + V \sum_{j,\alpha} [a_{j\alpha}^\dagger a_{j\alpha} - b_{j\alpha}^\dagger b_{j\alpha}] \\
&\quad + it_{\text{SO}} \sum_{j,\alpha,\beta} \sum_{k=1}^3 \sigma_z^{\alpha\beta} \left[[a_{j+\mu_k\alpha}^\dagger a_{j\beta} + b_{j\alpha}^\dagger b_{j+\mu_k\beta}] - [a_{j\alpha}^\dagger a_{j+\mu_k\beta} + b_{j+\mu_k\alpha}^\dagger b_{j\beta}] \right] \\
&\equiv H_{\text{KM}}^{\text{NN}} + H_{\text{KM}}^{\text{pot}} + H_{\text{KM}}^{\text{SO}},
\end{aligned} \tag{A.228}$$

where $j + \nu_k$ denotes the index of the k -th NN site of j and $j + \mu_k$ labels the index of the k -th NNN site in counterclockwise direction. Specifically, we have

$$\mathbf{R}_{j+\nu_k} = \mathbf{R}_j + \mathbf{a}_k \quad \text{and} \quad \mathbf{R}_{j+\mu_k} = \mathbf{R}_j - (-1)^k \mathbf{b}_k \equiv \mathbf{R}_j + \mathbf{b}'_k, \quad (\text{A.229})$$

as in Eq. (A.143). Apart from the extra spin index in the Kane–Mele model, we find that

$$H_{\text{KM}}^{\text{NN}} \simeq H_{\text{Hal}}^{\text{NN}} \quad \text{and} \quad H_{\text{KM}}^{\text{pot}} \simeq H_{\text{Hal}}^{\text{pot}} \quad (\text{A.230})$$

so we can adapt Eq. (A.144) to get

$$H_{\text{KM}}^{\text{NN}} = -t_{\text{hop}} \sum_{k=1}^3 \sum_{\mathbf{k}, \alpha} (a_{\mathbf{k}\alpha}^\dagger b_{\mathbf{k}\alpha}^\dagger) (\cos(\mathbf{k}\mathbf{a}_k) \tau_x - \sin(\mathbf{k}\mathbf{a}_k) \tau_y) \begin{pmatrix} a_{\mathbf{k}\alpha} \\ b_{\mathbf{k}\alpha} \end{pmatrix}, \quad (\text{A.231})$$

and Eq. (A.145) to get

$$H_{\text{KM}}^{\text{pot}} = V \sum_{\mathbf{k}, \alpha} (a_{\mathbf{k}\alpha}^\dagger b_{\mathbf{k}\alpha}^\dagger) \tau_z \begin{pmatrix} a_{\mathbf{k}\alpha} \\ b_{\mathbf{k}\alpha} \end{pmatrix}. \quad (\text{A.232})$$

Similarly, we may modify H_{NNN} from Eq. (A.146) to obtain

$$H_{\text{KM}}^{\text{SO}} = 2t_{\text{SO}} \sum_{k=1}^3 \sum_{\mathbf{k}, \alpha, \beta} \sigma_z^{\alpha\beta} (a_{\mathbf{k}\alpha}^\dagger b_{\mathbf{k}\alpha}^\dagger) \sin(\mathbf{k}\mathbf{b}'_k) \tau_z \begin{pmatrix} a_{\mathbf{k}\beta} \\ b_{\mathbf{k}\beta} \end{pmatrix}. \quad (\text{A.233})$$

Note once more that we write τ_j for the Pauli matrices associated with sublattice pseudospin and σ_j for the Pauli matrices describing electron spin. Combined, the total Fourier-transformed Kane–Mele Hamiltonian takes the form

$$H_{\text{KM}} = \sum_{\mathbf{k}} (a_{\mathbf{k}\uparrow}^\dagger b_{\mathbf{k}\uparrow}^\dagger a_{\mathbf{k}\downarrow}^\dagger b_{\mathbf{k}\downarrow}^\dagger) \left(h_0(\mathbf{k}) \sigma_z \otimes \tau_z + \mathbb{1}_2 \otimes [\mathbf{h}(\mathbf{k}) \boldsymbol{\tau}] \right) \begin{pmatrix} a_{\mathbf{k}\uparrow} \\ b_{\mathbf{k}\uparrow} \\ a_{\mathbf{k}\downarrow} \\ b_{\mathbf{k}\downarrow} \end{pmatrix} \quad (\text{A.234})$$

where

$$h_0(\mathbf{k}) = 2t_{\text{SO}} \sum_{j=1}^3 \sin(\mathbf{k}\mathbf{b}'_j) \quad (\text{A.235})$$

and

$$h_x(\mathbf{k}) = -t_{\text{hop}} \sum_{j=1}^3 \cos(\mathbf{k}\mathbf{a}_j) \quad , \quad h_y(\mathbf{k}) = t_{\text{hop}} \sum_{j=1}^3 \sin(\mathbf{k}\mathbf{a}_j) \quad , \quad h_z(\mathbf{k}) = V. \quad (\text{A.236})$$

Analogous to Eqs. (A.153) and (A.154 - A.156) we can use Eq. (A.150) to simplify Eqs. (A.235) and (A.236) as

$$h_0(\mathbf{k}) = -2t_{\text{SO}} [2 \sin(x) \cos(3y) - \sin(2x)] \quad (\text{A.237})$$

and

$$h_x(\mathbf{k}) = -t_{\text{hop}} [2 \cos(x) \cos(y) + \cos(2y)] \quad , \quad h_y(\mathbf{k}) = -t_{\text{hop}} [2 \cos(x) \sin(y) - \sin(2y)] \quad , \quad h_z(\mathbf{k}) = V. \quad (\text{A.238})$$

As the two spin projections in Eq. (A.234) do not mix we can diagonalise the spin-up and spin-down blocks separately. This yields the familiar energy dispersions

$$\begin{aligned} E_{\pm}^{\uparrow}(\mathbf{k}) &= \pm \sqrt{h_x(\mathbf{k})^2 + h_y(\mathbf{k})^2 + (h_z(\mathbf{k}) + h_0(\mathbf{k}))^2}, \\ E_{\pm}^{\downarrow}(\mathbf{k}) &= \pm \sqrt{h_x(\mathbf{k})^2 + h_y(\mathbf{k})^2 + (h_z(\mathbf{k}) - h_0(\mathbf{k}))^2}, \end{aligned} \quad (\text{A.239})$$

which may be combined as

$$E_{\pm}^{\sigma}(\mathbf{k}) = \pm \sqrt{h_x(\mathbf{k})^2 + h_y(\mathbf{k})^2 + (h_z(\mathbf{k}) + \eta_{\sigma} h_0(\mathbf{k}))^2} \quad (\text{A.240})$$

where $\eta_{\sigma} = \pm 1$ is positive (negative) for $\sigma = \uparrow$ ($\sigma = \downarrow$).

The Fourier transformed Kane–Mele Hamiltonian in Eq. (A.234) is given in terms of the four tensor basis matrices

$$\sigma_z \otimes \tau_z, \mathbb{1}_2 \otimes \tau_x, \mathbb{1}_2 \otimes \tau_y, \mathbb{1}_2 \otimes \tau_z. \quad (\text{A.241})$$

In this basis, the TRS operator takes the form $\mathcal{T} = (i\sigma_y \otimes \mathbb{1}_2)\mathcal{K}$ and the individual matrices transform as

$$\begin{aligned} \mathcal{T}(\sigma_z \otimes \tau_z) \mathcal{T}^\dagger &= (i\sigma_y \otimes \mathbb{1}_2)\mathcal{K}(\sigma_z \otimes \tau_z)\mathcal{K}((-i\sigma_y) \otimes \mathbb{1}_2) \\ &= (\sigma_y \otimes \mathbb{1}_2)(\sigma_z \otimes \tau_z)(\sigma_y \otimes \mathbb{1}_2) \\ &\stackrel{(\diamond)}{=} (\sigma_y \sigma_z \sigma_y \otimes \tau_z) \\ &\stackrel{(\diamond)}{=} (i\sigma_x \sigma_y \otimes \tau_z) \\ &= (i^2 \sigma_z \otimes \tau_z) \\ &= -(\sigma_z \otimes \tau_z) \\ \mathcal{T}(\mathbb{1}_2 \otimes \tau_x) \mathcal{T}^\dagger &= (i\sigma_y \otimes \mathbb{1}_2)\mathcal{K}(\mathbb{1}_2 \otimes \tau_x)\mathcal{K}((-i\sigma_y) \otimes \mathbb{1}_2) \\ &= (\sigma_y \otimes \mathbb{1}_2)(\mathbb{1}_2 \otimes \tau_x)(\sigma_y \otimes \mathbb{1}_2) \\ &= (\sigma_y^2 \otimes \tau_x) \\ &= (\mathbb{1}_2 \otimes \tau_x) \\ \mathcal{T}(\mathbb{1}_2 \otimes \tau_y) \mathcal{T}^\dagger &= (i\sigma_y \otimes \mathbb{1}_2)\mathcal{K}(\mathbb{1}_2 \otimes \tau_y)\mathcal{K}(-i\sigma_y \otimes \mathbb{1}_2) \\ &= (\sigma_y \otimes \mathbb{1}_2)(\mathbb{1}_2 \otimes (-\tau_y))(\sigma_y \otimes \mathbb{1}_2) \\ &= (\sigma_y^2 \otimes \tau_y) \\ &= -(\mathbb{1}_2 \otimes \tau_y) \\ \mathcal{T}(\mathbb{1}_2 \otimes \tau_z) \mathcal{T}^\dagger &= (i\sigma_y \otimes \mathbb{1}_2)\mathcal{K}(\mathbb{1}_2 \otimes \tau_z)\mathcal{K}((-i\sigma_y) \otimes \mathbb{1}_2) \\ &= (\sigma_y \otimes \mathbb{1}_2)(\mathbb{1}_2 \otimes \tau_z)(\sigma_y \otimes \mathbb{1}_2) \\ &= (\sigma_y^2 \otimes \tau_z) \\ &= (\mathbb{1}_2 \otimes \tau_z), \end{aligned} \quad (\text{A.242})$$

where we repeatedly used $\sigma_j \sigma_k = \delta_{kl} + i\epsilon_{jkl} \sigma_l$ in $(\diamond)$ and $\mathcal{K}\tau_y\mathcal{K} = -\tau_y$ while $\mathcal{K}\tau_x\mathcal{K} = \tau_x$ and $\mathcal{K}\tau_z\mathcal{K} = \tau_z$.

The $\mathbb{Z}_2$ Invariant of the Kane–Mele Model

Here we closely follow Ref. [181] to determine an explicit formula for the $\mathbb{Z}_2$ invariant of the Kane–Mele model. The TRS invariance of the Kane–Mele model means that we must consider “TRS-smooth” sections, i.e. sections $|u_n(k)\rangle$ fulfilling

$$|u_n(k)\rangle = \mathcal{T}|u_n(\theta(k))\rangle = \mathcal{T}|u_n(-k)\rangle, \quad (\text{A.243})$$

with the base space time-reversal involution $\theta : k \mapsto -k$. Just like the Chern number constitutes an obstruction to global sections, the $\mathbb{Z}_2$ invariant constitutes an obstruction to define global *TRS-smooth* sections. In order to construct TRS-smooth sections we introduce the so-called effective Brillouin zone $\mathcal{E}^2$, which is pictured in Fig. A.2. Generally, the effective Brillouin zone $\mathcal{E}^d$ is a fundamental domain of time-reversal symmetry, i.e. it contains exactly one representative from each pair k and $\theta(k) = -k$ of momenta related by TRS, up to points Γ_α that are invariant under TRS (time-reversal invariant momenta, or TRIMs). The boundary $\partial\mathcal{E}^d$ of $\mathcal{E}^d$ still supports a non-trivial action of TRS and we can again consider a fundamental domain $\mathcal{F}^{d-1}$ of $\partial\mathcal{E}^d$. In the present case, we have $d = 2$ and $\mathcal{E}^2$ is a cylinder $\mathbb{S}^1 \times [0, 1] \subset \mathbb{T}^2$ while $\mathcal{F}^1$ is the interval $[0, \pi]$. The concept of the effective Brillouin zone is crucial for the definition of the $\mathbb{Z}_2$ invariant.

Recall that the Kane–Mele Bloch bundle decomposes into a direct sum of two rank-two valence and conduction subbundles. At any TRIM Γ_α the time reversal operator $\mathcal{T}$ therefore acts as a skew-symmetric, anti-linear map on the rank-two fibre over Γ_α . In a local basis, this map is represented by a 2×2 skew-symmetric matrix. There are several natural quantities associated with such matrices, including their rank, determinant, and Pfaffian. The latter, though less commonly encountered, is particularly relevant

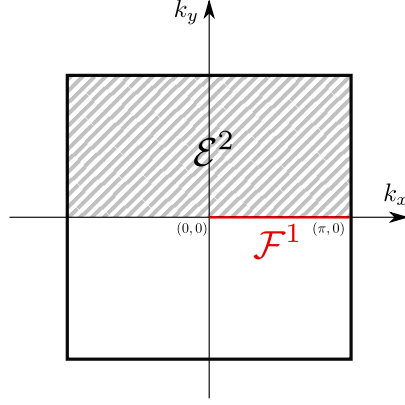


Figure A.2: Stylised two-dimensional effective Brillouin zone $\mathcal{E}^2$ (shaded) and the fundamental domain of its boundary $\mathcal{F}^1$ (red). Illustration created by the author, inspired by Ref. [181].

here. We have seen in Sec. 2.3.5 that the Pfaffian is an *invariant polynomial* and can therefore be used to generate topological invariants of a given vector bundle. For a 2×2 skew-symmetric matrix, the Pfaffian is especially simple: it equals the upper right entry.

$$\text{if } A = \begin{pmatrix} 0 & a \\ -a & 0 \end{pmatrix} \quad \text{then} \quad \text{Pf}(A) = a. \quad (\text{A.244})$$

Note that the actual choice of basis *does* matter because the sign of the Pfaffian changes for instance when two basis states are interchanged. The Kane–Mele invariant can be understood as a way of normalising and comparing these choices of Pfaffians at all four fixed points simultaneously. To see this, we utilise the fact that the Kramers pairs $\{|u_n^1(\Gamma_\alpha)\rangle, |u_n^2(\Gamma_\alpha)\rangle\}$ are linearly independent for any fixed point Γ_α so that we can choose them to satisfy

$$\mathcal{T}|u_n^1(\Gamma_\alpha)\rangle = |u_n^2(\Gamma_\alpha)\rangle \quad \text{and} \quad \mathcal{T}|u_n^2(\Gamma_\alpha)\rangle = -|u_n^1(\Gamma_\alpha)\rangle. \quad (\text{A.245})$$

We noted earlier that we would have to make sure that the sections $|u_n^j(k)\rangle$ are TRS-smooth, i.e. properly related between regions that are “time-reversed counterparts” of one another. To rigorously identify and connect these regions we make use of the notion of time-reversal stable, or θ -stable, regions. A subset $V \subset \mathbb{T}^2$ of the Brillouin torus is said to be stable under θ if it fulfils $\theta(V) = V$. Let us denote those θ -stable regions V that do not contain fixed points Γ_α by $\bar{V}$, i.e. $\bar{V} \cap \{\Gamma_\alpha\} = \emptyset$. Over every $\bar{V}$ we then have

$$|u_n^{1,2}(-k)\rangle = e^{i\chi_n^{1,2}(-k)} \mathcal{T}|u_n^{1,2}(k)\rangle \quad (\mathcal{T}^2 = -1) \quad e^{i\chi_n^{1,2}(-k)} = -e^{i\chi_n^{1,2}(k)}, \quad (\text{A.246})$$

with phase functions $\chi_n^{1,2}(k)$. One can show that these sections do not continuously extend to stable regions $\bar{V}$ that *do* contain fixed points, cf. Ref. [181]. However, it *is* possible to patch up a section using *both* $|u_n^1(k)\rangle$ and $|u_n^2(k)\rangle$ to bridge the TRIMs. To see this, we take some stable path P containing a fixed point Γ_α and decompose it as $P = P^- \amalg P^+$ where $\theta(V^-) = V^+$ and $V^- \cap V^+ = \Gamma_\alpha$. We then fix local sections over P^- and P^+ independently and continuously extend both sections to Γ_α from either side. Importantly, these combined sections can be chosen such that they *match* over Γ_α , such that the sections over P^- and P^+ can be “glued together” at Γ_α . In this way we can find two TRS-smooth sections that extend past Γ_α . Specifically, we denote the two combined sections by Roman numbers $|u_n^I(k)\rangle$ and $|u_n^{II}(k)\rangle$ and define

$$|u_n^I(k)\rangle = \begin{cases} |u_n^1(k)\rangle & \text{on } P^+ \\ |u_n^2(k)\rangle & \text{on } P^- \end{cases} \quad \text{and} \quad |u_n^{II}(k)\rangle = \begin{cases} |u_n^2(k)\rangle & \text{on } P^+ \\ |u_n^1(k)\rangle & \text{on } P^- \end{cases} \quad (\text{A.247})$$

where we glue together the functions $\chi_n^1(-k)$ and $\chi_n^2(k)$ at Γ_α to form a function $\chi_n(k)$. With this we get

$$|u_n^{II}(-k)\rangle = e^{i\chi_n(-k)} \mathcal{T}|u_n^I(k)\rangle \quad \text{and} \quad |u_n^I(-k)\rangle = -e^{i\chi_n(k)} \mathcal{T}|u_n^{II}(k)\rangle. \quad (\text{A.248})$$

In this basis, the $U(2)$ matrix representation of $\mathcal{T}$ reads

$$\omega_n(k) = (\langle u_n^s(-k) | \mathcal{T} | u_n^t(k) \rangle)_{s,t=I,II} = \begin{pmatrix} 0 & -e^{-i\chi_n(k)} \\ e^{-i\chi_n(-k)} & 0 \end{pmatrix}, \quad (\text{A.249})$$

which is skew-symmetric over the fixed point Γ_α . Its Pfaffian is $\text{Pf}(\omega_n(\Gamma_\alpha)) = -e^{-i\chi_n(\Gamma_\alpha)}$. From $|u_n^s(k)\rangle_{s=I,II}$ we may define Berry connections

$$\mathcal{A}_n^s(k) = i \langle u_n^s(k) | \partial_k | u_n^s(k) \rangle. \quad (\text{A.250})$$

With these we can compute so-called partial polarisations

$$P^s := \oint_C \mathcal{A}_n^s(k) dk \quad (\text{A.251})$$

over every closed curve $C \subset \mathbb{T}^2$. We can use the partial polarisations to define the *charge* polarisation $P_C = P^I + P^{II}$ and the *time-reversal* polarisation $P_T = P^I - P^{II}$. Now, the effective Brillouin zone boundary $\partial\mathcal{E}^2 = \{(k, 0) \mid k \in [-\pi, \pi]\} \subset \mathbb{T}^2$ corresponds to a closed curve $C \subset \mathbb{T}^2$, which is also a fundamental domain for the time-reversal involution θ . As a closed curve, $\partial\mathcal{E}^2$ contains two fixed points $\Gamma_1 = (0, 0)$ and $\Gamma_2 = (\pi, 0)$. Its time-reversal polarisation takes the form [181]

$$P_T = \frac{1}{2\pi} \int_{-\pi}^{\pi} (\mathcal{A}_n^I(k) - \mathcal{A}_n^{II}(k)) dk = \frac{1}{2\pi i} \left[\ln \frac{\det(\omega_n(\pi))}{\det(\omega_n(0))} - 2 \ln \frac{\text{Pf}(\omega_n(\pi))}{\text{Pf}(\omega_n(0))} \right]. \quad (\text{A.252})$$

Exponentiating both sides and using $e^{i\pi} = -1$ gives

$$(-1)^{P_T} = \prod_{\Gamma_\alpha = \Gamma_1, \Gamma_2} \frac{\sqrt{\det(\omega_n(\Gamma_\alpha))}}{\text{Pf}(\omega_n(\Gamma_\alpha))} = \prod_{\Gamma_\alpha = \Gamma_1, \Gamma_2} \text{sign}(\text{Pf}(\omega_n(\Gamma_\alpha))). \quad (\text{A.253})$$

The Kane–Mele invariant is then a mod two reduction

$$\nu := P_T \mod 2 \quad (\text{A.254})$$

of the time-reversal polarisation P_T .¹⁰ There are more formulas for the Kane–Mele invariant available in the literature, that highlight different properties. One wide-spread formula is

$$\nu := \frac{1}{2\pi} \left(\oint_{\partial\mathcal{E}^2} \mathcal{A} dk - \int_{\mathcal{E}^2} \mathcal{F} dk \right) \mod 2 \quad (\text{A.255})$$

with the Berry connection $\mathcal{A}$ and Berry curvature $\mathcal{F}$, see e.g. [181]. This formula is reminiscent of those yielding the Chern number as it solely depends on the Berry connection and curvature.

¹⁰It is important to point out that this is based on a continuous choice of the square root along the considered path.

A.11 Diagonalisation and Product Vacuum of BCS and Kitaev Chain Hamiltonians

Here, we provide a collection of explicit calculations for Chap. 9.

Bogoliubov Diagonalisation of Reduced BdG Hamiltonian

To begin with, we prove that the reduced BdG Hamiltonian Eq. (9.1) is diagonalised by the Bogoliubov diagonalisation Eq. (9.2). First, we show that the transformation matrix

$$U(\mathbf{k}) = \begin{pmatrix} u(\mathbf{k}) & -v(\mathbf{k}) \\ v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} \quad (\text{A.256})$$

is actually unitary. This is a direct consequence of Eq. (9.7), as

$$\begin{aligned} \begin{pmatrix} u(\mathbf{k}) & -v(\mathbf{k}) \\ v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} \begin{pmatrix} u(\mathbf{k}) & v(\mathbf{k}) \\ -v(\mathbf{k})^* & u(\mathbf{k}) \end{pmatrix} &= \begin{pmatrix} u(\mathbf{k})^2 + |v(\mathbf{k})|^2 & u(\mathbf{k})v(\mathbf{k}) - v(\mathbf{k})u(\mathbf{k}) \\ v(\mathbf{k})^*u(\mathbf{k}) - u(\mathbf{k})v(\mathbf{k})^* & |v(\mathbf{k})|^2 + u(\mathbf{k})^2 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}. \end{aligned} \quad (\text{A.257})$$

Second, $U(\mathbf{k})$ actually diagonalises the reduced BdG matrix as indicated which we will show by brute force in the following. For this, we drop the functional dependence on crystal momentum $\mathbf{k}$ to make the whole experience more readable and less painful. We get

$$\begin{aligned} U^\dagger \underline{h} U &= \begin{pmatrix} u & v \\ -v^* & u \end{pmatrix} \begin{pmatrix} \xi & \Delta \\ \Delta^* & -\xi \end{pmatrix} \begin{pmatrix} u & -v \\ v^* & u \end{pmatrix} \\ &= \begin{pmatrix} u & v \\ -v^* & u \end{pmatrix} \begin{pmatrix} \xi u + \Delta v^* & -\xi v + \Delta u \\ \Delta^* u - \xi v^* & -\Delta^* v - \xi u \end{pmatrix} \\ &= \begin{pmatrix} \xi u^2 + \Delta v^* u + v \Delta^* u - \xi |v|^2 & -\xi uv + \Delta u^2 - \Delta^* v^2 - \xi uv \\ -\xi uv^* - \Delta v^{*2} + \Delta^* u^2 - \xi v^* u & \xi |v|^2 - \Delta v^* u - v \Delta^* u - \xi u^2 \end{pmatrix} \\ &= \begin{pmatrix} A & B \\ B^* & -A \end{pmatrix}. \end{aligned} \quad (\text{A.258})$$

For A and B we find

$$\begin{aligned} A &= \xi u^2 + \Delta v^* u + \Delta^* v u - \xi |v|^2 \\ &= \xi \frac{1}{2} \left(1 + \frac{\xi}{E} \right) + \Delta \sqrt{\frac{1}{2} \left(1 - \frac{\xi}{E} \right)} \frac{\Delta^*}{|\Delta|} \sqrt{\frac{1}{2} \left(1 + \frac{\xi}{E} \right)} \\ &\quad + \Delta^* \sqrt{\frac{1}{2} \left(1 - \frac{\xi}{E} \right)} \frac{\Delta}{|\Delta|} \sqrt{\frac{1}{2} \left(1 + \frac{\xi}{E} \right)} - \xi \frac{1}{2} \left(1 - \frac{\xi}{E} \right) \frac{|\Delta|^2}{|\Delta|^2} \\ &= \frac{\xi}{2} \left(1 + \frac{\xi}{E} - 1 + \frac{\xi}{E} \right) + 2 \frac{|\Delta|^2}{|\Delta|} \sqrt{\frac{1}{2} \left(1 - \frac{\xi}{E} \right)} \frac{1}{2} \left(1 + \frac{\xi}{E} \right) \\ &= \frac{\xi^2}{E} + 2|\Delta| \sqrt{\frac{1}{4} \left(1 - \frac{\xi^2}{E^2} \right)} \\ &= \frac{\xi^2}{E} + |\Delta| \sqrt{\frac{E^2 - \xi^2}{E^2}} \\ &= \frac{\xi^2}{E} + |\Delta| \sqrt{\frac{|\Delta|^2}{E^2}} \\ &= \frac{\xi^2 + |\Delta|^2}{E} \\ &= E \end{aligned} \quad (\text{A.259})$$

and

$$\begin{aligned}
B &= -\xi uv + \Delta u^2 - \Delta^* v^2 - \xi uv \\
&= -2\xi uv + \Delta u^2 - \Delta^* v^2 \\
&= -2\xi \sqrt{\frac{1}{2} \left(1 + \frac{\xi}{E}\right)} \sqrt{\frac{1}{2} \left(1 - \frac{\xi}{E}\right)} \frac{\Delta}{|\Delta|} + \Delta \frac{1}{2} \left(1 + \frac{\xi}{E}\right) - \Delta^* \frac{1}{2} \left(1 - \frac{\xi}{E}\right) \frac{\Delta^2}{|\Delta|^2} \\
&= -2 \frac{\Delta \xi}{|\Delta|} \sqrt{\frac{1}{4} \left(1 + \frac{\xi}{E}\right) \left(1 - \frac{\xi}{E}\right)} + \frac{\Delta}{2} \left(1 + \frac{\xi}{E}\right) - \frac{\Delta}{2} \left(1 - \frac{\xi}{E}\right) \frac{|\Delta|^2}{|\Delta|^2} \\
&= -2 \frac{\Delta \xi}{|\Delta|} \sqrt{\frac{1}{4} \left(1 + \frac{\xi}{E}\right) \left(1 - \frac{\xi}{E}\right)} + \frac{\Delta}{2} \left(1 + \frac{\xi}{E}\right) - \frac{\Delta}{2} \left(1 - \frac{\xi}{E}\right) \\
&= -\frac{\Delta \xi}{|\Delta|} \sqrt{1 - \frac{\xi^2}{E^2}} + \frac{\Delta}{2} \left(1 + \frac{\xi}{E} - 1 + \frac{\xi}{E}\right) \\
&= -\frac{\Delta \xi}{|\Delta|} \sqrt{\frac{E^2 - \xi^2}{E^2}} + \frac{\Delta}{2} \frac{2\xi}{E} \\
&= -\frac{\Delta \xi}{|\Delta|} \sqrt{\frac{|\Delta|^2}{E^2}} + \frac{\xi \Delta}{E} \\
&= -\frac{\Delta \xi}{|\Delta|} \frac{|\Delta|}{E} + \frac{\xi \Delta}{E} \\
&= -\frac{\xi \Delta}{E} + \frac{\xi \Delta}{E} \\
&= 0.
\end{aligned} \tag{A.260}$$

so Eq. (A.258) becomes

$$U^\dagger \underline{h} U = \begin{pmatrix} u & v \\ -v^* & u \end{pmatrix} \begin{pmatrix} \xi & \Delta \\ \Delta^* & -\xi \end{pmatrix} \begin{pmatrix} u & -v \\ v^* & u \end{pmatrix} = \begin{pmatrix} A & B \\ B^* & -A \end{pmatrix} = \begin{pmatrix} E & 0 \\ 0 & -E \end{pmatrix}. \tag{A.261}$$

Fourier Transforming the BCS Chain Hamiltonian

Next, we explicitly Fourier transform the BCS tight-binding Hamiltonian from Eq. 9.8 using

$$c_{j\sigma} = \frac{1}{\sqrt{L}} \sum_{\mathbf{k}} e^{i\mathbf{k}\mathbf{R}_j} c_{\mathbf{k}\sigma} \tag{A.262}$$

to get

$$\begin{aligned}
H_{\text{BCS}} &= \sum_{j,\sigma} \left(-t \left[\frac{1}{L} \sum_{\mathbf{k},\mathbf{l}} e^{i(\mathbf{l}\mathbf{R}_{j+1} - \mathbf{k}\mathbf{R}_j)} c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{l}\sigma} + \frac{1}{L} \sum_{\mathbf{k},\mathbf{l}} e^{i(\mathbf{k}\mathbf{R}_j - \mathbf{l}\mathbf{R}_{j+1})} c_{\mathbf{l}\sigma}^\dagger c_{\mathbf{k}\sigma} \right] \right. \\
&\quad \left. - \mu \left[\frac{1}{L} \sum_{\mathbf{k},\mathbf{l}} e^{i(\mathbf{l}\mathbf{R}_j - \mathbf{k}\mathbf{R}_j)} c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{l}\sigma} \right] \right) \\
&\quad - \sum_j \left(\Delta \left[\frac{1}{L} \sum_{\mathbf{k},\mathbf{l}} e^{-i(\mathbf{l}\mathbf{R}_j + \mathbf{k}\mathbf{R}_j)} c_{\mathbf{k}\downarrow}^\dagger c_{\mathbf{l}\uparrow}^\dagger \right] + \Delta^* \left[\frac{1}{L} \sum_{\mathbf{k},\mathbf{l}} e^{i(\mathbf{l}\mathbf{R}_j + \mathbf{k}\mathbf{R}_j)} c_{\mathbf{k}\uparrow} c_{\mathbf{l}\downarrow} \right] \right) \\
&= \sum_{\mathbf{k},\mathbf{l},\sigma} \left(-t \left[\left(\frac{1}{L} \sum_j e^{i(\mathbf{l} - \mathbf{k})\mathbf{R}_j} \right) e^{i\mathbf{l}\mathbf{a}} c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{l}\sigma} + \left(\frac{1}{L} \sum_j e^{i(\mathbf{l} - \mathbf{k})\mathbf{R}_j} \right) e^{-i\mathbf{l}\mathbf{a}} c_{\mathbf{l}\sigma}^\dagger c_{\mathbf{k}\sigma} \right] \right. \\
&\quad \left. - \mu \left[\left(\frac{1}{L} \sum_j e^{i(\mathbf{l} - \mathbf{k})\mathbf{R}_j} \right) c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{l}\sigma} \right] \right) \\
&\quad - \sum_{\mathbf{k},\mathbf{l}} \left(\Delta \left[\left(\frac{1}{L} \sum_j e^{-i(\mathbf{l} + \mathbf{k})\mathbf{R}_j} \right) c_{\mathbf{k}\downarrow}^\dagger c_{\mathbf{l}\uparrow}^\dagger \right] + \Delta^* \left[\left(\frac{1}{L} \sum_j e^{i(\mathbf{l} + \mathbf{k})\mathbf{R}_j} \right) c_{\mathbf{k}\uparrow} c_{\mathbf{l}\downarrow} \right] \right), \tag{A.263}
\end{aligned}$$

which, using $\frac{1}{L} \sum_j e^{-i(l+\mathbf{k})\mathbf{R}_j} = \delta_{\mathbf{k}l}$ can be simplified to

$$\begin{aligned}
H_{\text{BCS}} &= \sum_{\mathbf{k}, \sigma} \left(-t \left[e^{i\mathbf{k}\mathbf{a}} c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{k}\sigma} + e^{-i\mathbf{k}\mathbf{a}} c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{k}\sigma} \right] - \mu c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{k}\sigma} \right) - \sum_{\mathbf{k}} \left(\Delta c_{\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}\uparrow}^\dagger + \Delta^* c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \right) \\
&= \sum_{\mathbf{k}, \sigma} \left(-t \left[e^{i\mathbf{k}\mathbf{a}} + e^{-i\mathbf{k}\mathbf{a}} \right] - \mu \right) c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{k}\sigma} - \sum_{\mathbf{k}, l} \left(\Delta c_{\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}\uparrow}^\dagger + \Delta^* c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \right) \\
&= \sum_{\mathbf{k}, l, \sigma} \left(-2t \cos(\mathbf{k}\mathbf{a}) - \mu \right) c_{\mathbf{k}\sigma}^\dagger c_{\mathbf{k}\sigma} - \sum_{\mathbf{k}} \left(\Delta c_{\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}\uparrow}^\dagger + \Delta^* c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \right) \\
&= \sum_{\mathbf{k}} \left(\xi(\mathbf{k}) (c_{\mathbf{k}\uparrow}^\dagger c_{\mathbf{k}\uparrow} + c_{-\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}\downarrow}) - \Delta c_{\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}\uparrow}^\dagger - \Delta^* c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} \right) \\
&= \sum_{\mathbf{k}} \left(\xi(\mathbf{k}) (c_{\mathbf{k}\uparrow}^\dagger c_{\mathbf{k}\uparrow} + 1 - c_{\mathbf{k}\downarrow} c_{\mathbf{k}\downarrow}^\dagger) + \Delta c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}^\dagger + \Delta^* c_{-\mathbf{k}\downarrow} c_{\mathbf{k}\uparrow} \right) \\
&= \sum_{\mathbf{k}} (c_{\mathbf{k}\uparrow}^\dagger, c_{-\mathbf{k}\downarrow}) \begin{pmatrix} \xi(\mathbf{k}) & \Delta \\ \Delta^* & -\xi(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}\uparrow}^\dagger \\ c_{-\mathbf{k}\downarrow} \end{pmatrix} + \sum_{\mathbf{k}} \xi(\mathbf{k}). \tag{A.264}
\end{aligned}$$

Note that we are considering a one-dimensional chain so that $\mathbf{k}\mathbf{a} = |\mathbf{k}||\mathbf{a}|$. Moreover, we usually set $a \equiv |\mathbf{a}| = 1$ so that $\mathbf{k}\mathbf{a} = |\mathbf{k}| = k$. As a result, we have

$$\xi(\mathbf{k}) = -2t \cos(\mathbf{k}\mathbf{a}) - \mu = -2t \cos(k) - \mu. \tag{A.265}$$

Product Form of the BCS Vacuum

Moving on, we briefly show that the BCS product vacuum does indeed take the form given in Eq. (9.15):

$$\begin{aligned}
|0\rangle_b &= \prod_{\mathbf{k}, \sigma} b_{\mathbf{k}\sigma} |0\rangle \\
&= \prod_{\mathbf{k}} \left(u(\mathbf{k}) c_{\mathbf{k}\uparrow} + v(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger \right) \left(-v(\mathbf{k}) c_{\mathbf{k}\uparrow}^\dagger + u(\mathbf{k}) c_{-\mathbf{k}\downarrow} \right) |0\rangle \\
&= \prod_{\mathbf{k}} \left(-u(\mathbf{k}) v(\mathbf{k}) c_{\mathbf{k}\uparrow}^\dagger c_{\mathbf{k}\uparrow} + \cancel{u(\mathbf{k})^2 c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}} - v(\mathbf{k})^2 c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger + \cancel{v(\mathbf{k}) u(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}\downarrow}} \right) |0\rangle \\
&= \prod_{\mathbf{k}} \left(-u(\mathbf{k}) v(\mathbf{k}) (1 - \cancel{c_{\mathbf{k}\uparrow}^\dagger c_{\mathbf{k}\uparrow}}) - v(\mathbf{k})^2 c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) |0\rangle \\
&= \prod_{\mathbf{k}} (-v(\mathbf{k})) \prod_{\mathbf{k}} (u(\mathbf{k}) + v_0(\mathbf{k}) e^{i\delta} c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle \\
&= \prod_{\mathbf{k}} (-v(\mathbf{k}) e^{i\delta}) \prod_{\mathbf{k}} (u(\mathbf{k}) e^{-i\delta} + v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle \tag{A.266}
\end{aligned}$$

where

$$\mathcal{N} := \prod_{\mathbf{k}} (-v(\mathbf{k}) e^{i\delta}) \tag{A.267}$$

is the norm since

$$\begin{aligned}
{}_b\langle 0|0\rangle_b &= \prod_{\mathbf{k}} |v(\mathbf{k})|^2 \left(\langle 0| \prod_{\mathbf{k}} (u(\mathbf{k}) e^{i\delta} + v_0(\mathbf{k}) c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow}) \right) \left(\prod_{\mathbf{k}} (u(\mathbf{k}) e^{-i\delta} + v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle \right) \\
&= |\mathcal{N}|^2 \langle 0| \prod_{\mathbf{k}} \left((u(\mathbf{k}) e^{i\delta} + v_0(\mathbf{k}) c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow}) (u(\mathbf{k}) e^{-i\delta} + v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) \right) |0\rangle \\
&= |\mathcal{N}|^2 \langle 0| \prod_{\mathbf{k}} \left(u(\mathbf{k})^2 + \cancel{u(\mathbf{k}) e^{i\delta} v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger} + v_0(\mathbf{k}) u(\mathbf{k}) e^{-i\delta} \cancel{c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow}} + v_0(\mathbf{k})^2 c_{\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow} c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) |0\rangle \\
&= |\mathcal{N}|^2 \prod_{\mathbf{k}} (u(\mathbf{k})^2 + v_0(\mathbf{k})^2) \\
&= |\mathcal{N}|^2. \tag{A.268}
\end{aligned}$$

Fourier Transforming the Kitaev Chain Hamiltonian

The Kitaev chain Hamiltonian from Eq. (9.26) Fourier transforms as

$$\begin{aligned}
H_K &= \sum_j \left(-t \left[\frac{1}{L} \sum_{\mathbf{k}, l} e^{i(l\mathbf{R}_j + 1 - \mathbf{k}\mathbf{R}_j)} c_{\mathbf{k}}^\dagger c_l + \frac{1}{L} \sum_{\mathbf{k}, l} e^{i(\mathbf{k}\mathbf{R}_j - l\mathbf{R}_{j+1})} c_l^\dagger c_{\mathbf{k}} \right] \right. \\
&\quad \left. - \mu \left[\frac{1}{L} \sum_{\mathbf{k}, l} e^{i(l\mathbf{R}_j - \mathbf{k}\mathbf{R}_j)} c_{\mathbf{k}}^\dagger c_l \right] \right. \\
&\quad \left. + \frac{\Delta}{2} \left[\frac{1}{L} \sum_{\mathbf{k}, l} e^{-i(\mathbf{k}\mathbf{R}_j + l\mathbf{R}_{j+1})} c_l^\dagger c_{\mathbf{k}}^\dagger - \frac{1}{L} \sum_{\mathbf{k}, l} e^{-i(\mathbf{k}\mathbf{R}_{j+1} + l\mathbf{R}_j)} c_l^\dagger c_{\mathbf{k}}^\dagger \right] \right. \\
&\quad \left. + \frac{\Delta^*}{2} \left[\frac{1}{L} \sum_{\mathbf{k}, l} e^{i(l\mathbf{R}_{j+1} + \mathbf{k}\mathbf{R}_j)} c_{\mathbf{k}} c_l - \frac{1}{L} \sum_{\mathbf{k}, l} e^{i(l\mathbf{R}_j + \mathbf{k}\mathbf{R}_{j+1})} c_{\mathbf{k}} c_l \right] \right) \\
&= \sum_{\mathbf{k}, l} \left(-t \left[\left(\frac{1}{L} \sum_j e^{i(l - \mathbf{k})\mathbf{R}_j} \right) e^{il\mathbf{a}} c_{\mathbf{k}}^\dagger c_l + \left(\frac{1}{L} \sum_j e^{i(l - \mathbf{k})\mathbf{R}_j} \right) e^{-il\mathbf{a}} c_l^\dagger c_{\mathbf{k}} \right] \right. \\
&\quad \left. - \mu \left[\left(\frac{1}{L} \sum_j e^{i(l - \mathbf{k})\mathbf{R}_j} \right) c_{\mathbf{k}}^\dagger c_l \right] \right. \\
&\quad \left. + \frac{\Delta}{2} \left[\left(\frac{1}{L} \sum_j e^{-i(l + \mathbf{k})\mathbf{R}_j} \right) e^{-il\mathbf{a}} c_l^\dagger c_{\mathbf{k}}^\dagger - \left(\frac{1}{L} \sum_j e^{-i(\mathbf{k} + l)\mathbf{R}_j} \right) e^{-i\mathbf{k}\mathbf{a}} c_l^\dagger c_{\mathbf{k}}^\dagger \right] \right. \\
&\quad \left. + \frac{\Delta^*}{2} \left[\left(\frac{1}{L} \sum_j e^{i(l + \mathbf{k})\mathbf{R}_j} \right) e^{il\mathbf{a}} c_{\mathbf{k}} c_l - \left(\frac{1}{L} \sum_j e^{i(\mathbf{k} + l)\mathbf{R}_j} \right) e^{i\mathbf{k}\mathbf{a}} c_{\mathbf{k}} c_l \right] \right) \\
&= \sum_{\mathbf{k}} \left([-t (e^{i\mathbf{k}\mathbf{a}} + e^{-i\mathbf{k}\mathbf{a}}) - \mu] c_{\mathbf{k}}^\dagger c_{\mathbf{k}} \right. \\
&\quad \left. + \frac{\Delta}{2} (e^{i\mathbf{k}\mathbf{a}} - e^{-i\mathbf{k}\mathbf{a}}) c_{-\mathbf{k}}^\dagger c_{\mathbf{k}}^\dagger + \frac{\Delta^*}{2} (e^{-i\mathbf{k}\mathbf{a}} - e^{i\mathbf{k}\mathbf{a}}) c_{\mathbf{k}} c_{-\mathbf{k}} \right) \\
&= \sum_{\mathbf{k}} \left([-2t \cos(\mathbf{k}\mathbf{a}) - \mu] c_{\mathbf{k}}^\dagger c_{\mathbf{k}} + i\Delta \sin(\mathbf{k}\mathbf{a}) c_{-\mathbf{k}}^\dagger c_{\mathbf{k}}^\dagger - i\Delta^* \sin(\mathbf{k}\mathbf{a}) c_{\mathbf{k}} c_{-\mathbf{k}} \right) \\
&\stackrel{(\diamond)}{=} \frac{1}{2} \sum_{\mathbf{k}} \left(\xi(\mathbf{k}) (c_{\mathbf{k}}^\dagger c_{\mathbf{k}} + c_{\mathbf{k}}^\dagger c_{\mathbf{k}}) - 2i\Delta \sin(\mathbf{k}\mathbf{a}) c_{\mathbf{k}}^\dagger c_{-\mathbf{k}}^\dagger + 2i\Delta^* \sin(\mathbf{k}\mathbf{a}) c_{-\mathbf{k}} c_{\mathbf{k}} \right) \\
&\stackrel{(\star)}{=} \frac{1}{2} \sum_{\mathbf{k}} \left(\xi(\mathbf{k}) (c_{\mathbf{k}}^\dagger c_{\mathbf{k}} + 1 - c_{\mathbf{k}} c_{\mathbf{k}}^\dagger) + \Delta(\mathbf{k}) c_{\mathbf{k}}^\dagger c_{-\mathbf{k}}^\dagger + \Delta(\mathbf{k})^* c_{-\mathbf{k}} c_{\mathbf{k}} \right) \\
&= \frac{1}{2} \sum_{\mathbf{k}} \left(\xi(\mathbf{k}) c_{\mathbf{k}}^\dagger c_{\mathbf{k}} - \xi(\mathbf{k}) c_{\mathbf{k}} c_{\mathbf{k}}^\dagger + \Delta(\mathbf{k}) c_{\mathbf{k}}^\dagger c_{-\mathbf{k}}^\dagger + \Delta(\mathbf{k})^* c_{-\mathbf{k}} c_{\mathbf{k}} \right) + \frac{1}{2} \sum_{\mathbf{k}} \xi(\mathbf{k}) \\
&= \frac{1}{2} \sum_{\mathbf{k}} \left(\xi(\mathbf{k}) c_{\mathbf{k}}^\dagger c_{\mathbf{k}} - \xi(-\mathbf{k}) c_{-\mathbf{k}} c_{-\mathbf{k}}^\dagger + \Delta(\mathbf{k}) c_{\mathbf{k}}^\dagger c_{-\mathbf{k}}^\dagger + \Delta(\mathbf{k})^* c_{-\mathbf{k}} c_{\mathbf{k}} \right) + E_0 \\
&= \frac{1}{2} \sum_{\mathbf{k}} \left(\xi(\mathbf{k}) c_{\mathbf{k}}^\dagger c_{\mathbf{k}} - \xi(\mathbf{k}) c_{-\mathbf{k}} c_{-\mathbf{k}}^\dagger + \Delta(\mathbf{k}) c_{\mathbf{k}}^\dagger c_{-\mathbf{k}}^\dagger + \Delta(\mathbf{k})^* c_{-\mathbf{k}} c_{\mathbf{k}} \right) + E_0 \\
&= \frac{1}{2} \sum_{\mathbf{k}} (c_{\mathbf{k}}^\dagger, c_{-\mathbf{k}}) \begin{pmatrix} \xi(\mathbf{k}) & \Delta(\mathbf{k}) \\ \Delta(\mathbf{k})^* & -\xi(\mathbf{k}) \end{pmatrix} \begin{pmatrix} c_{\mathbf{k}} \\ c_{-\mathbf{k}}^\dagger \end{pmatrix} + E_0, \tag{A.269}
\end{aligned}$$

where we defined $\xi(\mathbf{k}) = -2t \cos(\mathbf{k}\mathbf{a}) - \mu$ in $(\diamond)$, $\Delta(\mathbf{k}) = -2i\Delta \sin(\mathbf{k}\mathbf{a})$ in $(\star)$, and made repeated use of the fermionic anticommutator relations throughout.

Rearranging the Diagonalised Kitaev Chain Hamiltonian

The rearrangement of Eq. (9.30) into Eq. (9.34) reads

$$\begin{aligned}
H_K &= \frac{1}{2} \sum_{\mathbf{k}} (b_{\mathbf{k}}^\dagger, b_{-\mathbf{k}}) \begin{pmatrix} E(\mathbf{k}) & 0 \\ 0 & -E(\mathbf{k}) \end{pmatrix} \begin{pmatrix} b_{\mathbf{k}} \\ b_{-\mathbf{k}}^\dagger \end{pmatrix} \\
&= \frac{1}{2} \sum_{\mathbf{k}} E(\mathbf{k}) (b_{\mathbf{k}}^\dagger b_{\mathbf{k}} - b_{-\mathbf{k}} b_{-\mathbf{k}}^\dagger) \\
&= \frac{1}{2} \sum_{\mathbf{k}} E(\mathbf{k}) (b_{\mathbf{k}}^\dagger b_{\mathbf{k}} - (1 - b_{-\mathbf{k}}^\dagger b_{-\mathbf{k}})) \\
&= \frac{1}{2} \sum_{\mathbf{k}} E(\mathbf{k}) (b_{\mathbf{k}}^\dagger b_{\mathbf{k}} + b_{-\mathbf{k}}^\dagger b_{-\mathbf{k}}) + \sum_{\mathbf{k}} E(\mathbf{k}) \\
&= \frac{1}{2} \sum_{\mathbf{k}} E(\mathbf{k}) (b_{\mathbf{k}}^\dagger b_{\mathbf{k}} + b_{\mathbf{k}}^\dagger b_{\mathbf{k}}) + E'_0 \\
&= \sum_{\mathbf{k}} E(\mathbf{k}) b_{\mathbf{k}}^\dagger b_{\mathbf{k}} + E'_0
\end{aligned} \tag{A.270}$$

Product Form of the Kitaev Vacuum

The product vacuum of the Kitaev chain can be written as

$$\begin{aligned}
|0\rangle_b &= \prod_{\mathbf{0} < \mathbf{k} < \pi} b_{\mathbf{k}} b_{-\mathbf{k}} |0\rangle \\
&= \prod_{\mathbf{0} < \mathbf{k} < \pi} \left(u(\mathbf{k}) c_{\mathbf{k}\uparrow} + v(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger \right) \left(-v(\mathbf{k}) c_{\mathbf{k}\uparrow}^\dagger + u(\mathbf{k}) c_{-\mathbf{k}\downarrow} \right) |0\rangle \\
&= \prod_{\mathbf{0} < \mathbf{k} < \pi} \left(-u(\mathbf{k}) v(\mathbf{k}) c_{\mathbf{k}\uparrow} c_{\mathbf{k}\uparrow}^\dagger + \cancel{u(\mathbf{k})^2 c_{\mathbf{k}\uparrow}^\dagger c_{-\mathbf{k}\downarrow}} \xrightarrow{0} -v(\mathbf{k})^2 c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger + \cancel{v(\mathbf{k}) u(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{-\mathbf{k}\downarrow}} \xrightarrow{0} \right) |0\rangle \\
&= \prod_{\mathbf{0} < \mathbf{k} < \pi} \left(-u(\mathbf{k}) v(\mathbf{k}) (1 - \cancel{c_{\mathbf{k}\uparrow}^\dagger c_{\mathbf{k}\uparrow}}) \xrightarrow{0} -v(\mathbf{k})^2 c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) |0\rangle \\
&= \prod_{\mathbf{0} < \mathbf{k} < \pi} (-v(\mathbf{k})) \prod_{\mathbf{k}} (u(\mathbf{k}) + v(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle \\
&= \prod_{\mathbf{0} < \mathbf{k} < \pi} (-v(\mathbf{k}) e^{i\delta}) \prod_{\mathbf{k}} (u(\mathbf{k}) e^{-i\delta} + v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle
\end{aligned} \tag{A.271}$$

where

$$|\mathcal{N}| := \prod_{\mathbf{0} < \mathbf{k} < \pi} (-v(\mathbf{k}) e^{i\delta}) \tag{A.272}$$

is once again the norm because

$$\begin{aligned}
{}_d \langle 0|0 \rangle_d &= \prod_{\mathbf{0} < \mathbf{k} < \pi} |v(\mathbf{k})|^2 \left(\langle 0| \prod_{\mathbf{0} < \mathbf{k}' < \pi} (u(\mathbf{k}') e^{i\delta} + v_0(\mathbf{k}') c_{\mathbf{k}'\downarrow} c_{-\mathbf{k}'\uparrow}) \right) \left(\prod_{\mathbf{0} < \mathbf{k} < \pi} (u(\mathbf{k}) e^{-i\delta} + v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) |0\rangle \right) \\
&= |\mathcal{N}|^2 \langle 0| \prod_{\mathbf{0} < \mathbf{k} < \pi} \left((u(\mathbf{k}) e^{i\delta} + v_0(\mathbf{k}) c_{\mathbf{k}\downarrow} c_{-\mathbf{k}\uparrow}) (u(\mathbf{k}) e^{-i\delta} + v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger) \right) |0\rangle \\
&= |\mathcal{N}|^2 \langle 0| \prod_{\mathbf{0} < \mathbf{k} < \pi} \left(u(\mathbf{k})^2 + \cancel{u(\mathbf{k}) e^{i\delta} v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow} c_{\mathbf{k}\uparrow}^\dagger} \xrightarrow{0} + \cancel{v_0(\mathbf{k}) u(\mathbf{k}) e^{-i\delta} c_{\mathbf{k}\downarrow} c_{-\mathbf{k}\uparrow}} \xrightarrow{0} + v_0(\mathbf{k})^2 c_{\mathbf{k}\downarrow} c_{-\mathbf{k}\uparrow} c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger \right) |0\rangle \\
&= |\mathcal{N}|^2 \prod_{\mathbf{0} < \mathbf{k} < \pi} (u(\mathbf{k})^2 + v_0(\mathbf{k})^2) \\
&= |\mathcal{N}|^2.
\end{aligned} \tag{A.273}$$

One normalised Kitaev ground state is therefore

$$|0\rangle_d = \prod_{\mathbf{0} < \mathbf{k} < \pi} (u(\mathbf{k}) e^{-i\delta} + v_0(\mathbf{k}) c_{-\mathbf{k}\downarrow}^\dagger c_{\mathbf{k}\uparrow}^\dagger). \tag{A.274}$$

A.12 Technical Aspects of the Kitaev Chain Model

Here, we provide a collection of explicit calculations for Chap. 10.

Majorana Representation of the Kitaev Chain

Following Sec. 5.5, we may introduce self-adjoint Majorana operators

$$\gamma_j^A = c_j^\dagger + c_j = \gamma_j^{A\dagger} \quad \text{and} \quad \gamma_j^B = i(c_j^\dagger - c_j) = \gamma_j^{A\dagger}, \quad (\text{A.275})$$

which give

$$c_j = \frac{1}{2}(\gamma_j^A + i\gamma_j^B) \quad \text{and} \quad c_j^\dagger = \frac{1}{2}(\gamma_j^A - i\gamma_j^B). \quad (\text{A.276})$$

If we plug these into Eq. (10.1) we get

$$\begin{aligned} H_{\text{Kit}} &= \frac{1}{4} \left(\sum_{j=1}^{L-1} \left[-t ((\gamma_{j+1}^A - i\gamma_{j+1}^B)(\gamma_j^A + i\gamma_j^B) + (\gamma_j^A - i\gamma_j^B)(\gamma_{j+1}^A + i\gamma_{j+1}^B)) \right. \right. \\ &\quad \left. \left. + \Delta_0 (\gamma_j^A + i\gamma_j^B)(\gamma_{j+1}^A + i\gamma_{j+1}^B) + \Delta_0 (\gamma_{j+1}^A - i\gamma_{j+1}^B)(\gamma_j^A - i\gamma_j^B) \right] \right. \\ &\quad \left. - \mu \sum_{j=1}^L ((\gamma_j^A - i\gamma_j^B)(\gamma_j^A + i\gamma_j^B) - 2) \right) \\ &= \frac{1}{4} \left(\sum_{j=1}^{L-1} \left[-t (\cancel{\gamma_{j+1}^A \gamma_j^A} - i\gamma_{j+1}^B \gamma_j^A + i\gamma_{j+1}^A \gamma_j^B + \cancel{\gamma_{j+1}^B \gamma_j^B} + \cancel{\gamma_j^A \gamma_{j+1}^A} + i\gamma_j^A \gamma_{j+1}^B - i\gamma_j^B \gamma_{j+1}^A + \cancel{\gamma_j^B \gamma_{j+1}^B}) \right. \right. \\ &\quad \left. \left. + \Delta_0 (\cancel{\gamma_j^A \gamma_{j+1}^A} + i\gamma_j^B \gamma_{j+1}^A + i\gamma_j^A \gamma_{j+1}^B + \cancel{\gamma_j^B \gamma_{j+1}^B}) \right. \right. \\ &\quad \left. \left. + \Delta_0 (\cancel{\gamma_{j+1}^A \gamma_j^A} - i\gamma_{j+1}^B \gamma_j^B - i\gamma_{j+1}^A \gamma_j^A + \cancel{\gamma_{j+1}^B \gamma_j^B}) \right] \right. \\ &\quad \left. - \mu \sum_{j=1}^L (\gamma_j^A \gamma_j^A + i\gamma_j^A \gamma_j^B - i\gamma_j^B \gamma_j^A + \gamma_j^B \gamma_j^B - 2) \right) \\ &= \frac{1}{4} \left(\sum_{j=1}^{L-1} \left[-2it (\gamma_j^A \gamma_{j+1}^B - i\gamma_j^B \gamma_{j+1}^A) + 2i\Delta_0 (\gamma_j^B \gamma_{j+1}^A + \gamma_j^A \gamma_{j+1}^B) \right] - \mu \sum_{j=1}^L (2i\gamma_j^A \gamma_j^B + \mathbb{Z} - \mathbb{Z}) \right) \\ &= \frac{i}{2} \left(\sum_{j=1}^{L-1} \left[(\Delta_0 + t) \gamma_j^B \gamma_{j+1}^A + (\Delta_0 - t) \gamma_j^A \gamma_{j+1}^B \right] - \mu \sum_{j=1}^L \gamma_j^A \gamma_j^B \right), \quad (\text{A.277}) \end{aligned}$$

where we have repeatedly used $\{\gamma_j^X, \gamma_k^Y\} = 2\delta_{jk}\delta_{XY}$.

Energy Gap in the Kitaev Chain Model

From

$$\begin{aligned} \frac{\partial}{\partial k} |\mathbf{h}(\mathbf{k})|^2 &= \frac{\partial}{\partial k} (4\Delta_0^2 \sin^2(k) + (\mu + 2t \cos(k))^2) \\ &= 8\Delta_0^2 \cos(k) \sin(k) - 4t \sin(k) (\mu + 2t \cos(k)) \\ &= 4 \sin(k) [2\Delta_0^2 \cos(k) - t(\mu + 2t \cos(k))] \\ &\stackrel{!}{=} 0 \end{aligned} \quad (\text{A.278})$$

we get two types of extrema, namely

$$4 \sin(k) \stackrel{!}{=} 0 \quad \implies \quad k_{c,1}(n) = n\pi \quad \text{with} \quad n \in \mathbb{Z}, \quad (\text{A.279})$$

which exists for all values of t, μ and Δ_0 , and

$$2\Delta_0^2 \cos(k) - t(\mu + 2t \cos(k)) \stackrel{!}{=} 0 \quad \implies \quad k_{c,2} = \pm \arccos \left(\frac{t\mu}{2(\Delta_0^2 - t^2)} \right), \quad (\text{A.280})$$

which only exists if

$$\left| \frac{t\mu}{2(\Delta_0^2 - t^2)} \right| \leq 1. \quad (\text{A.281})$$

With

$$\begin{aligned} |\mathbf{h}(\mathbf{k}_{c,1}(n))|^2 &= 4\Delta_0^2 \sin^2(k_{c,1}) + (\mu + 2t \cos(k_{c,1}))^2 \\ &= (\mu + (-1)^n 2t)^2 \end{aligned} \quad (\text{A.282})$$

and

$$\begin{aligned} |\mathbf{h}(\mathbf{k}_{c,2})|^2 &= 4\Delta_0^2 \sin^2(k_{c,2}) + (\mu + 2t \cos(k_{c,2}))^2 \\ &= 4\Delta_0^2 (1 - \cos^2(k_{c,2})) + (\mu + 2t \cos(k_{c,2}))^2 \\ &= 4\Delta_0^2 \left(1 - \left| \frac{t\mu}{2(\Delta_0^2 - t^2)} \right|^2 \right) + \left(\mu + 2t \left[\frac{t\mu}{2(\Delta_0^2 - t^2)} \right] \right)^2 \\ &= \frac{16\Delta_0^2(\Delta_0^2 - t^2)^2 - 4\Delta_0^2 t^2 \mu^2 + (2\mu(\Delta_0^2 - t^2) + 2t^2 \mu)^2}{4(\Delta_0^2 - t^2)^2} \\ &= \frac{16\Delta_0^2(\Delta_0^2 - t^2)^2 - 4\Delta_0^2 t^2 \mu^2 + 4\mu^2 \Delta_0^4}{4(\Delta_0^2 - t^2)^2} \\ &= \frac{16\Delta_0^2(\Delta_0^2 - t^2)^2 + 4\Delta_0^2 \mu^2 (\Delta_0^2 - t^2)}{4(\Delta_0^2 - t^2)^2} \\ &= \frac{4\Delta_0^2(\Delta_0^2 - t^2) + \Delta_0^2 \mu^2}{(\Delta_0^2 - t^2)} \\ &= \Delta_0^2 \left(4 + \frac{\mu^2}{(\Delta_0^2 - t^2)} \right) \end{aligned} \quad (\text{A.283})$$

we then get a band gap of

$$\Delta E = 2 \cdot \min \left(|\mu + 2t|, |\mu - 2t|, |\Delta_0| \sqrt{4 + \frac{\mu^2}{(\Delta_0^2 - t^2)}} \right). \quad (\text{A.284})$$

Depending on which of these is realised, the band gap is situated at

$$\Delta E = \begin{cases} 2|\mu + 2t| & \text{at } k = 0 \\ 2|\mu - 2t| & \text{at } k = \pm\pi \\ 2|\Delta_0| \sqrt{4 + \frac{\mu^2}{(\Delta_0^2 - t^2)}} & \text{at } k = \pm \arccos \left(\frac{t\mu}{2(\Delta_0^2 - t^2)} \right) \end{cases}. \quad (\text{A.285})$$

For $\Delta_0 > 0$ we get band closures at $k = 0$ ($k = \pm\pi$) when $\mu = -2t$ ($\mu = 2t$). If $\Delta_0 = 0$, we additionally get band closures at

$$k = \pm \arccos \left(-\frac{\mu}{2t} \right) \quad (\text{A.286})$$

when $|\mu| \leq 2|t|$.

Degree Invariant of the Kitaev Chain Model

To compute the degree of $m : \mathbb{S}_k^1 \rightarrow \mathbb{S}_m^1$ from Eq. (10.38), we first need to determine a normalised volume form on the target manifold $\mathbb{S}_m^1 \subset \mathbb{R}^3$. To this end, we observe that $m(\mathbf{k})$ traces out a circle lying in the two-dimensional plane

$$P = \text{span} \left\{ \begin{pmatrix} \sin(\phi) \\ \cos(\phi) \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right\} \subset \mathbb{R}^3. \quad (\text{A.287})$$

We may therefore set $\phi = 0$, noting that (i) $|\mathbf{h}(\mathbf{k})|$ is independent¹¹ of ϕ , and (ii) the vector

$$\mathbf{h}(\mathbf{k}) = \begin{pmatrix} 2\Delta_0 \sin(\phi) \sin(k) \\ -2\Delta_0 \cos(\phi) \sin(k) \\ \mu + 2t \cos(k) \end{pmatrix} \quad (\text{A.288})$$

from Eqs. (10.20) and (10.21) itself induces a homotopy between $m(\mathbf{k})$ at $\phi = 0$ and $m(\mathbf{k})$ at any other value of ϕ . Taken together, (i) and (ii) ensure that the map m remains well-defined and that its degree $\deg m$ remains invariant under variations of ϕ . For $\phi = 0$ we get

$$m : \mathbb{S}_{\mathbf{k}}^1 \rightarrow \mathbb{S}_m^1 \subset \mathbb{R}^2 \quad , \quad \mathbf{k} \mapsto \frac{\mathbf{h}(\mathbf{k})}{|\mathbf{h}(\mathbf{k})|} \quad , \quad (\text{A.289})$$

with

$$\mathbf{h}(\mathbf{k}) = \begin{pmatrix} -2\Delta_0 \sin(k) \\ \mu + 2t \cos(k) \end{pmatrix} . \quad (\text{A.290})$$

This allows us to limit our calculation to the target manifold $\mathbb{S}_m^1 \subset \mathbb{R}^2$. The canonical normalised volume form of $\mathbb{S}^1 \subset \mathbb{R}^2$ reads

$$\eta = \frac{1}{2\pi} (x_1 dx_2 - x_2 dx_1) = \frac{1}{2\pi} \epsilon^{ab} x_a dx_b \quad (\text{A.291})$$

where we introduced the totally antisymmetric tensor $\epsilon^{ab} = -\epsilon^{ba}$ for convenience. Now for the pullback. The pullback of a differential form by some map is something us physicists tend to be rather wary of. Let me break down the rough pullback recipe before we actually plunge into business and use it. Let X be a one-dimensional space and let by Y a two-dimensional space like in our case. Given a smooth map $m : X \rightarrow Y$ and a one form $\omega \in \Omega^1(Y)$ then the pullback $m^*(\omega)$ of ω by m is a one form of X , i.e. $m^*\omega \in \Omega^1(X)$. We call m^* a pullback map with respect to m because it points the opposite direction of m : while m goes from X to Y the pullback m^* points from $\Omega^1(Y)$ to $\Omega^1(X)$. A one form $\omega \in \Omega^1(Y)$ evaluated at $y = (y_1, y_2) \in Y$ reads

$$\omega[y] = \omega_1(y) dy_1 + \omega_2(y) dy_2 \quad (\text{A.292})$$

where $(\omega_1(y), \omega_2(y)) \in \mathbb{R}^2$ and $dy_1, dy_2 \in T_y(Y)^*$ form the dual basis to the basis $\{\partial/\partial y_1, \partial/\partial y_2\}$ of $T_y(Y)$. Recall that $\omega[y] \in T_y(Y)^*$ eats one vector $v \in T_y(Y)$ (a tangent vector at $y \in Y$) as

$$\begin{aligned} \omega[y](v) &= \omega[y] \left(v^1 \frac{\partial}{\partial y_1} + v^2 \frac{\partial}{\partial y_2} \right) \\ &= \omega_1(y) v^1 + \omega_2(y) v^2 \end{aligned} \quad (\text{A.293})$$

which is a real number, i.e. $\omega[y](v) \in \mathbb{R}$. The pullback is going to be our solution to the problem of associating a one form $m^*\omega \in \Omega^1(X)$ to any given $\omega \in \Omega^1(Y)$ by our map $m : X \rightarrow Y$. So what could be a natural course of action facing this? Given a point $x \in X$ and a vector $u \in T_x(X)$, the natural idea is to use the point $y = (m_1(x), m_2(x)) \in Y$ and to push forward the vector u as $m_*(u) \in T_{m(x)}(Y)$. This is exactly how the pullback is defined, i.e.

$$(m^*\omega)[x](u) := \omega[m(x)](m_*(u)). \quad (\text{A.294})$$

Now we “only” have to compute the push forward $m_*(u)$ which is a much easier problem. It is defined as

$$m_*(u) = dm_1[m(x)](u) \frac{\partial}{\partial y_1} + dm_2[m(x)](u) \frac{\partial}{\partial y_2} \quad (\text{A.295})$$

which yields

$$(m^*\omega)[x](u) = \omega_1(m(x)) dy_1 dm_1[m(x)](u) \frac{\partial}{\partial y_1} + \omega_2(m(x)) dy_2 dm_2[m(x)](u) \frac{\partial}{\partial y_2}. \quad (\text{A.296})$$

¹¹ As $|\mathbf{h}(\mathbf{k})| = \sqrt{4\Delta_0^2 \sin^2(k)(\sin^2(\phi) + \cos^2(\phi)) + (\mu + 2t \cos(k))^2} = \sqrt{4\Delta_0^2 \sin^2(k) + (\mu + 2t \cos(k))^2}$.

By definition, the terms $dy_i \partial/\partial y_i$ cancel out and we get

$$(m^* \omega)[x](u) = \omega_1(m(x)) dm_1[m(x)](u) + \omega_2(m(x)) dm_2[m(x)](u) \quad (\text{A.297})$$

and since we are interested in what happens for the one form itself we can drop the test argument $u \in T_x(X)$ and end up with

$$(m^* \omega)[x] = \omega_1(m(x)) dm_1[m(x)] + \omega_2(m(x)) dm_2[m(x)]. \quad (\text{A.298})$$

In our case we have the mapping $m : k \mapsto (m_1(k), m_2(k))$ and the normalised volume form $\eta = \frac{1}{2\pi}(-x_2, x_1)$ of $\mathbb{S}^1 \subset \mathbb{R}^2$ which is given in the $\{dx_1, dx_2\}$ basis of $T_x(\mathbb{S}^1 \subset \mathbb{R}^2)^*$. Let us compute the individual parts of Eq. (A.298). Namely, we get

$$(\eta_1[m(k)], \eta_2[m(k)]) = \left(-\frac{1}{2\pi} m_2(k), \frac{1}{2\pi} m_1(k)\right) \quad (\text{A.299})$$

as well as

$$dm_1[m(x)] = \frac{\partial m_1(k)}{\partial k} dk \quad \text{and} \quad dm_2[m(x)] = \frac{\partial m_2(k)}{\partial k} dk \quad (\text{A.300})$$

which gives

$$\begin{aligned} (m^* \eta)[k] &= -\frac{1}{2\pi} m_2(k) \frac{\partial m_1(k)}{\partial k} dk + \frac{1}{2\pi} m_1(k) \frac{\partial m_2(k)}{\partial k} dk \\ &= \frac{1}{2\pi} \epsilon^{ab} m_a(k) \frac{\partial m_b(k)}{\partial k} dk. \end{aligned} \quad (\text{A.301})$$

The degree of m is then the integral

$$\begin{aligned} \deg m &= \int_{\mathbb{S}^1} m^* \eta \\ &= \frac{1}{2\pi} \int_{\mathbb{S}^1 \subset \mathbb{R}^2} \epsilon^{ab} m_a(k) \frac{\partial m_b(k)}{\partial k} dk \end{aligned} \quad (\text{A.302})$$

In order to efficiently compute this integral we will use a trick: we will rewrite Eq. (A.302) as a complex integral and use Cauchy's argument principle for meromorphic¹² functions. To do this consider

$$\epsilon^{ab} m_a(k) \frac{\partial m_b(k)}{\partial k} = m_1(k) \frac{\partial m_2(k)}{\partial k} - m_2(k) \frac{\partial m_1(k)}{\partial k}. \quad (\text{A.303})$$

With Eq. (A.289) we get

$$\begin{aligned} m_1 \frac{\partial m_2}{\partial k} &= \frac{h_1}{\sqrt{h_1^2 + h_2^2}} \left(\frac{h'_2}{\sqrt{h_1^2 + h_2^2}} + h_2 \frac{\partial}{\partial k} \left[\frac{1}{\sqrt{h_1^2 + h_2^2}} \right] \right) \\ m_2 \frac{\partial m_1}{\partial k} &= \frac{h_2}{\sqrt{h_1^2 + h_2^2}} \left(\frac{h'_1}{\sqrt{h_1^2 + h_2^2}} + h_1 \frac{\partial}{\partial k} \left[\frac{1}{\sqrt{h_1^2 + h_2^2}} \right] \right). \end{aligned} \quad (\text{A.304})$$

Here we dropped the k -dependence for better readability. We will stick to that convention for most of the remainder of the paragraph. When taking the difference between both terms in Eq. (A.304) to get to Eq. (A.303) the terms with derivatives on $1/\sqrt{h_1^2 + h_2^2}$ cancel out and we get

$$\epsilon^{ab} m_a(k) \frac{\partial m_b(k)}{\partial k} = \frac{h_1 h'_2 - h_2 h'_1}{h_1^2 + h_2^2}. \quad (\text{A.305})$$

Note that denominator equals the squared energy dispersion from Eq. (10.24) as

$$h_1^2 + h_2^2 = 4\Delta_0^2 \sin^2(k) + (\mu + 2t \cos(k))^2 \equiv E(k)^2. \quad (\text{A.306})$$

¹²A meromorphic function f on an open set U is a function that is holomorphic on all of U except for a set of isolated points $\{p\}$ which are poles of f . A holomorphic function is a complex valued function that is complex differentiable in an open neighbourhood of each point in its complex domain.

This already tells us that the degree of m is only going to be well-defined for as long as the bulk spectrum remains gapped. In order to enter the realm of complex analysis we convince ourselves that the following equality holds:

$$\frac{\partial}{\partial k} [(h_1 \pm ih_2)] (h_1 \mp ih_2) = h'_1 h_1 + h_2 h'_2 \mp i(h'_1 h_2 - h'_2 h_1) \quad (\text{A.307})$$

where the derivative with respect to k acts on the first factor only and where the $\pm$ and $\mp$ signs signify that there is an ambiguity in our description. Namely, the notation is supposed to show that there are two versions of the equation: we could either choose plus in the first factor and minus in the second factor or minus in the first factor and plus in the second. Above, we claimed that this choice is but a description ambiguity so we should be able to use either version of Eq. (A.307) and obtain the same result eventually. We will put up with the notational overhead in order to show that this is, indeed, the case. Comparing Eq. (A.307) to Eq. (A.305) we find

$$\epsilon^{ab} m_a(k) \frac{\partial m_b(k)}{\partial k} = \pm i \left[\frac{\partial}{\partial k} [(h_1 \pm ih_2)] (h_1 \mp ih_2) - h'_1 h_1 - h_2 h'_2 \right] \quad (\text{A.308})$$

which gives

$$\deg m = \pm \frac{i}{2\pi} \int_{\mathbb{S}_k^1} \frac{\frac{\partial}{\partial k} [(h_1 \pm ih_2)] (h_1 \mp ih_2) - (h'_1 h_1 + h_2 h'_2)}{h_1^2 + h_2^2} dk. \quad (\text{A.309})$$

The second term is easily evaluated as

$$\mp \frac{i}{2\pi} \int_{\mathbb{S}_k^1} \frac{(h'_1 h_1 + h_2 h'_2)}{h_1^2 + h_2^2} dk = \mp \frac{i}{4\pi} [\ln(h_1^2 + h_2^2)]_{k=0}^{2\pi} = 0 \quad (\text{A.310})$$

because the bulk dispersion $h_1^2 + h_2^2 = E(k)^2$ is 2π periodic. This leaves us with

$$\begin{aligned} \deg m &= \pm \frac{i}{2\pi} \int_{\mathbb{S}_k^1} \frac{\frac{\partial}{\partial k} [(h_1 \pm ih_2)] (h_1 \mp ih_2)}{h_1^2 + h_2^2} dk \\ &= \pm \frac{i}{2\pi} \int_{\mathbb{S}_k^1} \frac{\frac{\partial}{\partial k} [(h_1 \pm ih_2)] (h_1 \mp ih_2)}{(h_1 \pm ih_2)(h_1 \mp ih_2)} dk \\ &= \pm \frac{i}{2\pi} \int_{\mathbb{S}_k^1} \frac{\frac{\partial}{\partial k} [(h_1 \pm ih_2)]}{(h_1 \pm ih_2)} dk. \end{aligned} \quad (\text{A.311})$$

where in the second line we added a bit of trivial notation for the denominator to clarify that the spare $\mp$ terms cancel out to yield the last line of the equation. Let us reintroduce the k -dependence and define the complex valued function

$$\mathcal{M}^\pm(k) := h_1(k) \pm ih_2(k) \quad (\text{A.312})$$

and use it to evaluate Eq. (A.311) as

$$\begin{aligned} \deg m &= \pm \frac{i}{2\pi} \int_{\mathbb{S}_k^1} \frac{\frac{\partial}{\partial k} \mathcal{M}^\pm(k)}{\mathcal{M}^\pm(k)} dk \\ &= \pm \frac{i}{2\pi} \int_{\mathbb{S}_k^1} \frac{\partial}{\partial k} \ln \mathcal{M}^\pm(k) dk \end{aligned} \quad (\text{A.313})$$

This expression for the degree of m is a complex contour integral of the form

$$\frac{i}{2\pi} \oint_C \frac{f'(z)}{f(z)} dz \quad (\text{A.314})$$

where $f(z)$ is any meromorphic function on $C \cup C^\circ$ and where C is a simple¹³ contour. The sweet thing about complex contour integrals of the type in Eq. (A.314) is that for them Cauchy's argument principle applies. It states that

$$\frac{i}{2\pi} \oint_C \frac{f'(z)}{f(z)} dz = P_f - Z_f \quad (\text{A.315})$$

¹³A simple contour is a counterclockwise oriented contour without self-intersections. In our case, we deal with the contour $\mathbb{S}^1 \subset \mathbb{C} \sim \{z \in \mathbb{C} \mid |z| = 1\}$ which is the stereotypical example of a simple contour.

where P_f and Z_f denote the number of poles and zeros of $f(z)$ inside the contour C with each pole and zero counted as many times as its order and multiplicity indicate. Using Cauchy's argument principle on Eq. (A.313) we get

$$\begin{aligned}\deg m &= \pm \frac{i}{2\pi} \int_{\mathbb{S}_k^1} \frac{\partial}{\partial k} \ln \mathcal{M}^\pm(k) dk \\ &= \pm \frac{i}{2\pi} \oint_C \frac{\partial}{\partial z} \ln \mathcal{M}^\pm(z) dz \\ &= \pm (P_{\mathcal{M}^\pm} - Z_{\mathcal{M}^\pm})\end{aligned}\tag{A.316}$$

where we embedded $\mathbb{S}_k^1 \hookrightarrow \mathbb{C}$ as $k \mapsto z = e^{ik}$ to parameterise the contour

$$C := \{z \in \mathbb{C} : |z| = 1\} \subset \mathbb{C}\tag{A.317}$$

and where $P_{\mathcal{M}}$ and $Z_{\mathcal{M}}$ denote the number of poles and zeros of $\mathcal{M}^\pm(z)$ inside the complex unit circle. We can now write down our function $\mathcal{M}^\pm(z)$ as

$$\begin{aligned}\mathcal{M}^\pm(k) &= h_1(k) \pm i h_2(k) \\ &= -2\Delta_0 \sin k \pm i(\mu + 2t \cos k) \\ &= -2\Delta_0 \frac{1}{2i} (e^{ik} - e^{-ik}) \pm i(\mu + 2t \frac{1}{2} (e^{ik} + e^{-ik})) \\ &\stackrel{z=e^{ik}}{=} i\Delta_0 \left(z - \frac{1}{z}\right) \pm i(\mu + t \left(z + \frac{1}{z}\right)) \\ &= i \left[(\Delta_0 \pm t)z - (\Delta_0 \mp t) \frac{1}{z} \pm \mu \right]\end{aligned}\tag{A.318}$$

and examine its poles and zeros for its analytic continuation to the embedding $z \in \mathbb{C}_{|z| \leq 1}$ of the unit disc into the complex plane. Clearly, the zeros and poles of $\mathcal{M}^\pm(z)$ are dependent on the parameters t, Δ_0 and μ . In order to explicitly determine them, we first take

$$\mathcal{M}^+(z) = i \left[(\Delta_0 + t)z - (\Delta_0 - t) \frac{1}{z} + \mu \right]\tag{A.319}$$

and find that it has a pole at $z = 0$ iff $t \neq \Delta_0$ and that it has no zeros iff $t = -\Delta_0$. For $t \neq -\Delta_0$ we can determine the zeros via

$$\begin{aligned}0 &= i \left[(\Delta_0 \pm t)z - (\Delta_0 \mp t) \frac{1}{z} \pm \mu \right] \Leftrightarrow 0 = (\Delta_0 + t)z^2 - (\Delta_0 - t) + \mu z \\ &= z^2 - \frac{(\Delta_0 - t)}{(\Delta_0 + t)} + \frac{\mu}{(\Delta_0 + t)} z\end{aligned}\tag{A.320}$$

which gives

$$z_\pm = \frac{-\mu \pm \sqrt{\mu^2 + 4(\Delta_0^2 - t^2)}}{2(\Delta_0 + t)}.\tag{A.321}$$

For $\mu = 0$ we can simplify this further and get

$$\begin{aligned}z_\pm &= \pm \sqrt{\frac{(\Delta_0^2 - t^2)}{(\Delta_0 + t)^2}} \\ &= \pm \sqrt{\frac{(\Delta_0 - t)}{(\Delta_0 + t)}}.\end{aligned}\tag{A.322}$$

With this we can compute $\deg m$ for the two special points associated to phases (A) and (B)

$$\begin{aligned}(A) : \quad \Delta_0 = t = 0, \mu \neq 0 &\implies P_{\mathcal{M}^+} = 0, Z_{\mathcal{M}^+} = 0 &\implies \deg m = 0 \\ (B) : \quad \Delta_0 = t \neq 0, \mu = 0 &\implies P_{\mathcal{M}^+} = 0, Z_{\mathcal{M}^+} = 1 &\implies \deg m = -1.\end{aligned}\tag{A.323}$$

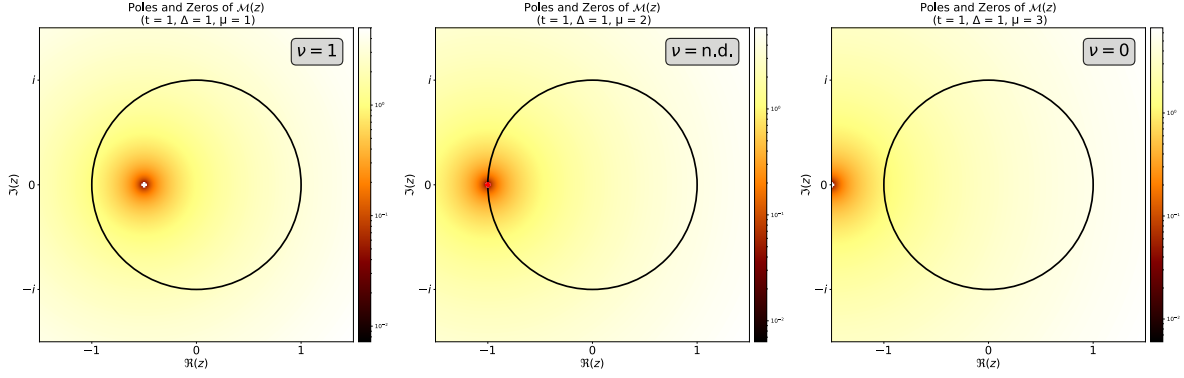


Figure A.3: Location of the zero (white plus) of the analytic continuation of $\mathcal{M}(z) := \mathcal{M}^+(z)$ in the square region $[-1, 1] \times [-i, i] \subset \mathbb{C}$. Parameters are listed in the top of each panel. Critical poles and zeros that are located on the unit circle are indicated by red markers of the respective type. The values of the degree and topological invariant of the Kitaev chain are shown in the top right corner. The contour of integration, the unit circle, is shown in black and the background heatmap shows the absolute $|\mathcal{M}(z)|$ on the unit square.

If we take

$$\mathcal{M}^-(z) = i \left[(\Delta_0 - t)z - (\Delta_0 + t)\frac{1}{z} - \mu \right] \quad (\text{A.324})$$

instead of $\mathcal{M}^+(z)$ we find that it has a pole at $z = 0$ iff $t \neq -\Delta_0$ and that it has no zeros iff $t = \Delta_0$. For $t \neq \Delta_0$ we can determine the zeros in the same fashion as before obtaining

$$z_{\pm} = \frac{\mu \pm \sqrt{\mu^2 + 4(\Delta_0^2 - t^2)}}{2(\Delta_0 - t)} \quad (\text{A.325})$$

which for $\mu = 0$ becomes

$$z_{\pm} = \pm \sqrt{\frac{(\Delta_0 + t)}{(\Delta_0 - t)}}. \quad (\text{A.326})$$

Using this we get

$$\begin{aligned} (A): \quad \Delta_0 = t = 0, \mu \neq 0 &\implies P_{\mathcal{M}^-} = 0, Z_{\mathcal{M}^-} = 0 &\implies \deg m = 0 \\ (B): \quad \Delta_0 = t \neq 0, \mu = 0 &\implies P_{\mathcal{M}^+} = 1, Z_{\mathcal{M}^-} = 0 &\implies \deg m = -1. \end{aligned} \quad (\text{A.327})$$

for the two special points associated to phases (A) and (B). This is the same result that we got when we used $\mathcal{M}^+(z)$ instead of $\mathcal{M}^-(z)$, namely

$$\begin{aligned} (A): \quad \Delta_0 = t = 0, \mu \neq 0 &\implies P_{\mathcal{M}^{\pm}} = 0, Z_{\mathcal{M}^{\pm}} = 0 &\implies \deg m = 0 \\ (B): \quad \Delta_0 = t \neq 0, \mu = 0 &\implies \begin{cases} P_{\mathcal{M}^+} = 0, Z_{\mathcal{M}^+} = 1 \\ P_{\mathcal{M}^-} = 1, Z_{\mathcal{M}^-} = 0 \end{cases} &\implies \deg m = -1. \end{aligned} \quad (\text{A.328})$$

which shows that the degree is in fact independent of the arbitrary sign we introduced earlier. The above considerations identify phase (A) to be topologically trivial and phase (B) to be topologically non-trivial. Note that our embedding $\mathbb{S}^1 \hookrightarrow \mathbb{C}$ determines the global sign of the degree. For example, using $k \mapsto z^* = e^{-ik}$ clearly reverses the unit circle direction and, indeed, we find that

$$\mathcal{M}^{\pm}(z^*) = i \left[(\Delta_0 \pm t)\frac{1}{z^*} - (\Delta_0 \mp t)z^* \pm \mu \right] \quad (\text{A.329})$$

produces an overall sign yielding $\deg m = +1$ in the topologically non-trivial region. This is expected because the embedding is an inherent part of the map m itself. Different embeddings define different

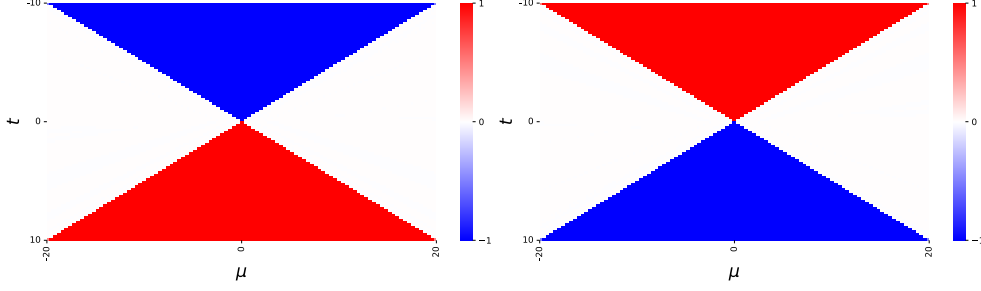


Figure A.4: Diagram of $\deg m$ in the (t, μ) plane for $\Delta_0 = +1$ (left) and $\Delta_0 = -1$ (right).

maps with (potentially) different degrees.¹⁴ In case of the Kitaev chain, the presence of an antilinear particle hole symmetry (PHS) induces a relation $z \leftrightarrow z^*$ between a complex number and its complex conjugate so the two solutions $\deg m = +1$ and $\deg m = -1$ correspond to the same topological phase. The topological index ν of the physical Kitaev model is therefore not determined by the degree $\deg m$ of the map m itself, but rather by its modulus $|\deg m|$. If we continuously move away from the special phase (A) and (B) points in the parameter space

$$\mathcal{P} = \text{span} \{t, \mu, \Delta_0\}, \quad (\text{A.330})$$

we effectively push around the locations of zeros or poles of $\mathcal{M}^\pm(z)$ in the interior $C^\circ = \{z \in \mathbb{C} : |z| < 1\}$ of C . According to Cauchy’s argument principle, this cannot change the value of the integral, and hence the degree $\deg m$ of m . Conversely, the only way to change the value of $\deg m$ is to tune parameters such that a zero or a pole leaves the region enclosed by our contour C . However, the only way to *continuously* push a zero or a pole out of that region is to make it cross the contour C at some point. This is shown in Fig. A.3. Recall that the absolute value of our complex valued function $\mathcal{M}^\pm(z)$ from Eq. (A.312) equals the absolute value of the bulk energy dispersion $E(k = i \ln z)$ for all z on the contour $C = \{z \in \mathbb{C} : |z| = 1\}$. A parameter configuration $p \in \mathcal{P}$ for which there exist $z \in C$ where $|\mathcal{M}^\pm(z)| = 0$, like the one in the middle panel of Fig. A.3, therefore signifies a band closure of the physical model. At the same time, the degree becomes ill-defined at such a parameter configuration because Cauchy’s argument principle forbids the function to have zeros or poles on the contour of integration. Fig. A.4 shows the degree of the map $m(k)$ in the (t, μ) plane for $\Delta_0 = +1$ and $\Delta_0 = -1$. As expected, we find $\deg m = 0$ for $2|t| < |\mu|$ and $\deg m = \pm 1$ for $2|t| > |\mu|$. Any two regions with different $\deg m$ are separated by a gap closure either due to $|\mu| = 2|t|$ or $\Delta_0 = 0$, cf. Eq. (A.285). There are two noteworthy subtleties to consider here. The first one is that there are parameter regions, such as $2t > |\mu|$ at $\Delta_0 = +1$ and $-2t > |\mu|$ at $\Delta_0 = -1$, that have the same degree $\deg m = +1$, but are still discontinuously separated by an energy gap closure. At first glance, this seems to conflict with the *completeness* of the degree as a homotopy invariant, which ensures that any two maps $f, g : \mathbb{S}^n \rightarrow \mathbb{S}^n$ are homotopic *if and only if* $\deg f = \deg g$. The resolution lies in the fact that our maps $m(k) : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ are far from arbitrary: they arise from specific parameter-dependent Hamiltonians and are therefore fully determined by points p in our parameter space $\mathcal{P}$. The existence of discontinuously (by gap closure) separated parameter regions with the same degree does not contradict the completeness of the degree as a homotopy invariant because it only precludes the existence of a continuous path within the highly specialised family of Hamiltonian maps $m(k)$. In particular, the matching degree still guarantees that $f_0(k) = m_{p_0}(k)$ and $f_1(k) = m_{p_1}(k)$ are homotopic as abstract maps $f : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ – just not through any homotopy confined to the “physical” subset of Hamiltonian maps.

The second subtlety in Fig. A.4 is that $\deg m \mapsto -\deg m$ under $\Delta \mapsto -\Delta$. This appears to contradict the fact that $m_{\phi=0}(k)$ for $\phi = 0$ is homotopic to $m_\phi(k)$ for any other value of ϕ , which should render the sign of the superconducting gap Δ irrelevant, since $\Delta_{\phi=0} = -\Delta_{\phi=\pi}$ while $\deg m_{\phi=0}(k) = \deg m_{\phi=\pi}(k)$. To resolve this, we note that, despite being closely related, the “two-dimensional” mapping $m : \mathbb{S}_k^1 \rightarrow \mathbb{S}_m^1 \subset \mathbb{R}^2$ from Eq. (A.289) and the “three-dimensional” mapping $m : \mathbb{S}_k^1 \rightarrow \mathbb{S}_m^1 \subset \mathbb{R}^3$ from Eq. (10.38) define distinct maps. In particular, either map is based on a specific choice of orientations for $\mathbb{S}_k^1$ and $\mathbb{S}^1 \subset \mathbb{R}^n$,

¹⁴In contrast, the choice of whether we use $\mathcal{M}^\pm$ to compute the integral is a technical one that does not effect the result as is seen in Eq. (A.328).

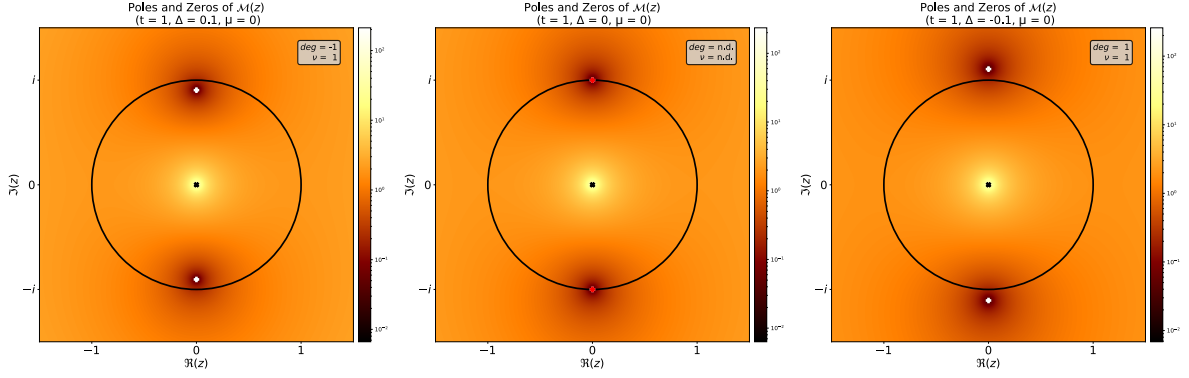


Figure A.5: Location of pole (black cross marker) and zeros (white plus markers) of the analytic continuation of $\mathcal{M}(z) := \mathcal{M}^+(z)$ in the square region $[-1, 1] \times [-i, i] \subset \mathbb{C}$. Parameters are listed in the top of each panel. Critical poles and zeros that are located on the unit circle are indicated by red markers of the respective type. The values of the degree and topological invariant of the Kitaev chain are shown in the top right corner. The contour of integration, the unit circle, is shown in black and the background heatmap shows the absolute $|\mathcal{M}(z)|$ on the unit square.

which has to be carefully accounted for in any comparison between them. To see this, we point out that at any arbitrary but fixed value of ϕ , the associated three-dimensional mapping $m_\phi : \mathbb{S}_k^1 \rightarrow \mathbb{S}_m^1 \subset \mathbb{R}^3$ can be restricted to a two-dimensional mapping $m : \mathbb{S}_k^1 \rightarrow \mathbb{S}_m^1 \subset \mathbb{R}^2$. Specifically, each three-dimensional map is actually a map

$$m_\phi : \mathbb{S}_k^1 \rightarrow \mathbb{S}_m^1 \subset P_\phi \subset \mathbb{R}^3 \quad (\text{A.331})$$

from $\mathbb{S}_k^1$ to the ϕ -dependent plane $P_\phi \subset \mathbb{R}^3$ from Eq. (A.287). Using that P_ϕ is readily isomorphic to $\mathbb{R}^2$, we then get a two-dimensional map

$$m_\phi : \mathbb{S}_k^1 \rightarrow \mathbb{S}_m^1 \subset \mathbb{R}^2, \quad (\text{A.332})$$

where the ϕ subscript indicates that the orientation of $\mathbb{S}_m^1 \subset P_\phi \subset \mathbb{R}^3$ is taken into account. The subtlety arises precisely for pairs $(\phi, \phi + \pi)$ where $P_\phi = P_{\phi+\pi}$, but with opposite orientations. This is easily seen considering P from Eq. (A.287) for $\phi = 0$ and $\phi = \pi$,

$$P_0 = \text{span} \left\{ \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right\} \quad \text{and} \quad P_\pi = \text{span} \left\{ \begin{pmatrix} 0 \\ -1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right\}, \quad (\text{A.333})$$

whose canonic orientations

$$\mathbf{O}_\phi = \mathbf{v}_{\phi,1} \times \mathbf{v}_{\phi,2} \quad (\text{A.334})$$

as subspaces of $\mathbb{R}^3$ are

$$\mathbf{O}_0 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \times \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \quad \text{and} \quad \mathbf{O}_\pi = \begin{pmatrix} 0 \\ -1 \\ 0 \end{pmatrix} \times \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} -1 \\ 0 \\ 0 \end{pmatrix}, \quad (\text{A.335})$$

readily satisfying

$$\mathbf{O}_0 = -\mathbf{O}_\pi. \quad (\text{A.336})$$

The relative sign flip of the degree between $\Delta_0 = +1$ and $\Delta_0 = -1$ in Fig. A.4 does not contradict the homotopy equivalence between $m_{\phi=0}(k)$ and $m_{\phi=\pi}(k)$ because it is undone by taking into account the reversal of orientation between $m_{\phi=0}(k)$ and $m_{\phi=\pi}(k)$. This is illustrated by Fig. A.5, which shows how the two zeros of the Kitaev map first lie inside the unit circle contour C for $\Delta_0 = 0.1$, then cross C at $z = \pm i$ at $\Delta_0 = 0$ and eventually lie outside C at $\Delta_0 = -0.1$. Throughout this transition, the degree changes from $\deg m = -1$ ($\nu = 1$) over $\deg m = \nu = \text{n.d.}$ to $\deg m = 1$ ($\nu = 1$) in the process. The relative sign between $\deg m = -1$ at $\Delta_0 = 0.1$ and $\deg m = +1$ at $\Delta_0 = -0.1$ can be undone by reversing the orientation of the integration contour, cf. Eq. (A.315).

The Form of the Topological Majorana Boundary Modes

Here we motivate the form Eq. (10.44) of the topological Majorana boundary modes in the Kitaev chain. Following [264], we start with the Majorana representation of H_{Kit} from Eq. (A.277). For now we assume the translationally invariant limit $L \rightarrow \infty$ such that we can Fourier transform the Majorana field operators as

$$\gamma_j^{A/B} = \frac{1}{\sqrt{L}} \sum_{\mathbf{k}} e^{-i\mathbf{k}\mathbf{R}_j} \gamma_{\mathbf{k}}^{A/B}. \quad (\text{A.337})$$

Note that due to the Majorana reality condition $\gamma_j = \gamma_j^\dagger$ in real space we get

$$\gamma_{\mathbf{k}}^{A/B\dagger} = \frac{1}{\sqrt{L}} \sum_{j=1}^L e^{-i\mathbf{k}\mathbf{R}_j} \gamma_j^{A/B\dagger} = \frac{1}{\sqrt{L}} \sum_{j=1}^L e^{-i\mathbf{k}\mathbf{R}_j} \gamma_j^{A/B} = \gamma_{-\mathbf{k}}^{A/B} \quad (\text{A.338})$$

in (quasi-)momentum representation. With this the Kitaev chain Hamiltonian in Eq. (10.4) becomes

$$\begin{aligned} H_{\text{Kit}} &= \frac{i}{2} \left(\sum_{j=1}^{L-1} \left[(\Delta_0 + t) \gamma_j^B \gamma_{j+1}^A + (\Delta_0 - t) \gamma_j^A \gamma_{j+1}^B \right] - \mu \sum_{j=1}^L \gamma_j^A \gamma_j^B \right) \\ &= \frac{i}{2L} \sum_{j,\mathbf{k},\mathbf{k}'} \left[-\mu e^{-i(\mathbf{k}+\mathbf{k}')\mathbf{R}_j} \gamma_{\mathbf{k}}^A \gamma_{\mathbf{k}'}^B + (\Delta_0 + t) e^{i(\mathbf{k}+\mathbf{k}')\mathbf{R}_j} e^{i\mathbf{k}'\mathbf{a}} \gamma_{\mathbf{k}}^B \gamma_{\mathbf{k}'}^A + (\Delta_0 - t) e^{i(\mathbf{k}+\mathbf{k}')\mathbf{R}_j} e^{i\mathbf{k}'\mathbf{a}} \gamma_{\mathbf{k}}^A \gamma_{\mathbf{k}'}^B \right] \\ &\stackrel{(\odot)}{=} \frac{i}{2} \sum_{\mathbf{k}} \left[-\mu \gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B + (\Delta_0 + t) e^{-i\mathbf{k}\mathbf{a}} \gamma_{\mathbf{k}}^B \gamma_{-\mathbf{k}}^A + (\Delta_0 - t) e^{-i\mathbf{k}\mathbf{a}} \gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B \right] \\ &= \frac{i}{2} \sum_{\mathbf{k}} \left[-\frac{\mu}{2} (\gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B + \gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B) \right. \\ &\quad \left. + \frac{(\Delta_0 + t)}{2} e^{-i\mathbf{k}\mathbf{a}} (\gamma_{\mathbf{k}}^B \gamma_{-\mathbf{k}}^A + \gamma_{\mathbf{k}}^B \gamma_{-\mathbf{k}}^A) + \frac{(\Delta_0 - t)}{2} e^{-i\mathbf{k}\mathbf{a}} (\gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B + \gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B) \right] \\ &= \frac{i}{2} \sum_{\mathbf{k}} \left[-\frac{\mu}{2} (\gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B + \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right. \\ &\quad \left. + \frac{(\Delta_0 + t)}{2} (e^{-i\mathbf{k}\mathbf{a}} \gamma_{\mathbf{k}}^B \gamma_{-\mathbf{k}}^A + e^{i\mathbf{k}\mathbf{a}} \gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A) + \frac{(\Delta_0 - t)}{2} (e^{-i\mathbf{k}\mathbf{a}} \gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B + e^{i\mathbf{k}\mathbf{a}} \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right] \\ &= \frac{i}{4} \sum_{\mathbf{k}} \left[\mu (\gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A - \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right. \\ &\quad \left. + \Delta_0 (e^{-i\mathbf{k}\mathbf{a}} \gamma_{\mathbf{k}}^B \gamma_{-\mathbf{k}}^A + e^{i\mathbf{k}\mathbf{a}} \gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A + e^{-i\mathbf{k}\mathbf{a}} \gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B + e^{i\mathbf{k}\mathbf{a}} \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right. \\ &\quad \left. + t (e^{-i\mathbf{k}\mathbf{a}} \gamma_{\mathbf{k}}^B \gamma_{-\mathbf{k}}^A + e^{i\mathbf{k}\mathbf{a}} \gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A - e^{-i\mathbf{k}\mathbf{a}} \gamma_{\mathbf{k}}^A \gamma_{-\mathbf{k}}^B - e^{i\mathbf{k}\mathbf{a}} \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right] \\ &= \frac{i}{4} \sum_{\mathbf{k}} \left[\mu (\gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A - \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right. \\ &\quad \left. + \Delta_0 ([e^{i\mathbf{k}\mathbf{a}} - e^{-i\mathbf{k}\mathbf{a}}] \gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A + [e^{i\mathbf{k}\mathbf{a}} - e^{-i\mathbf{k}\mathbf{a}}] \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right. \\ &\quad \left. + t ([e^{i\mathbf{k}\mathbf{a}} + e^{-i\mathbf{k}\mathbf{a}}] \gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A - [e^{i\mathbf{k}\mathbf{a}} + e^{-i\mathbf{k}\mathbf{a}}] \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right] \\ &= \frac{i}{4} \sum_{\mathbf{k}} \left[\mu (\gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A - \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right. \\ &\quad \left. + \Delta_0 (2i \sin(\mathbf{k}\mathbf{a}) \gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A + 2i \sin(\mathbf{k}\mathbf{a}) \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) + t (2 \cos(\mathbf{k}\mathbf{a}) \gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A - 2 \cos(\mathbf{k}\mathbf{a}) \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right] \\ &= \frac{i}{4} \sum_{\mathbf{k}} \left[\mu (\gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A - \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) + 2i\Delta_0 \sin(\mathbf{k}\mathbf{a}) (\gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A + \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) + 2t \cos(\mathbf{k}\mathbf{a}) (\gamma_{-\mathbf{k}}^B \gamma_{\mathbf{k}}^A - \gamma_{-\mathbf{k}}^A \gamma_{\mathbf{k}}^B) \right] \\ &\stackrel{(\star)}{=} \frac{i}{4} \sum_{\mathbf{k}} (\gamma_{-\mathbf{k}}^A, \gamma_{-\mathbf{k}}^B) h_{\text{Kit}}(\mathbf{k}) \begin{pmatrix} \gamma_{\mathbf{k}}^A \\ \gamma_{\mathbf{k}}^B \end{pmatrix} \\ &\equiv \frac{i}{4} \sum_{\mathbf{k}} \phi^\dagger(\mathbf{k}) h_{\text{Kit}}(\mathbf{k}) \phi(\mathbf{k}), \quad (\text{A.339}) \end{aligned}$$

where we repeatedly used the general delta summation identity

$$\frac{1}{L} \sum_j e^{i(\mathbf{k}-\mathbf{k}')\mathbf{R}_j} = \delta(\mathbf{k}-\mathbf{k}') \quad (\text{A.340})$$

in $(\diamond)$, and deployed the anticommutation relation of Majorana operators along with a renaming of $\mathbf{k} \mapsto -\mathbf{k}$ in some parts of the expression to arrive at a form where we may define a 2×2 BdG matrix

$$h_{\text{Kit}}(\mathbf{k}) = -2\Delta_0 \sin(\mathbf{k}\mathbf{a}) \sigma_x + (\mu + 2t \cos(\mathbf{k}\mathbf{a})) \sigma_y \quad (\text{A.341})$$

in $(\star)$. Note that because of Eq. (A.338) we can indeed understand Eq. (A.339) in the usual BdG sense. Now, we are interested in exponentially localised MZMs of H_{Kit} that appear when we consider large but finite instances of the topological Kitaev chain. Following [264], we can find such localised modes by looking for *non-oscillatory* solutions $\Psi^{A/B}(\mathbf{k})$ of

$$h_{\text{Kit}}(\mathbf{k}) \Psi^{A/B}(\mathbf{k}) = 0 \quad (\text{A.342})$$

where A/B refers to the two Majorana species in the system. To probe for such solutions, we simply replace the plane-wave factor $e^{i\mathbf{k}\mathbf{a}}$ by

$$e^{i\mathbf{k}\mathbf{a}} \mapsto e^{-\mathbf{k}\mathbf{a}} =: e^{-q}. \quad (\text{A.343})$$

This means that we are effectively attempting to diagonalise a Hamiltonian without translation invariance using eigenstates of the translation-invariant problem. As we will see shortly, this *is* possible but only if suitable boundary conditions can be satisfied. Namely, we have to ensure that the wave function cannot leave the lattice.

If we substitute Eq. (A.343) into Eq. (A.341) we get

$$h_{\text{Kit}}(q) = 2i\Delta_0 \sinh(q) \sigma_x + (\mu + 2t \cosh(q)) \sigma_y \quad (\text{A.344})$$

the zero solutions of which are easiest obtained by squaring the BdG matrix $h_{\text{Kit}}(q)$ first. Due to the algebra of the Pauli matrices we get

$$\begin{aligned} h_{\text{BdG}}(q)^2 &= -4\Delta_0^2 \sinh^2(q) \sigma_x^2 - 2i\Delta_0 \sinh(q) (\mu + 2t \cosh(q)) (\sigma_x \sigma_y + \sigma_y \sigma_x) + (\mu + 2t \cosh(q))^2 \sigma_y^2 \\ &= \left(-4\Delta_0^2 \sinh^2(q) + (\mu + 2t \cosh(q))^2 \right) \mathbb{1} \end{aligned} \quad (\text{A.345})$$

so we need to solve

$$0 \stackrel{!}{=} -4\Delta_0^2 \sinh^2(q) + (\mu + 2t \cosh(q))^2, \quad (\text{A.346})$$

which, after rearranging and taking the square root, reads

$$(\mu + 2t \cosh(q)) \stackrel{!}{=} \pm 2\Delta_0 \sinh(q). \quad (\text{A.347})$$

If we choose the positive solution and spell out the hyperbolicals we get

$$\mu + t(e^q + e^{-q}) \stackrel{!}{=} \Delta_0(e^q - e^{-q}). \quad (\text{A.348})$$

Now we may solve this equation for either e^q or e^{-q} . We settle for e^{-q} and get

$$\begin{aligned} 0 &\stackrel{!}{=} (t + \Delta_0)e^{-2q} + \mu e^{-q} + (t - \Delta_0) \\ &= e^{-2q} + \frac{\mu}{(t + \Delta_0)} e^{-q} + \frac{(t - \Delta_0)}{(t + \Delta_0)} \\ &=: x^2 + \frac{\mu}{(t + \Delta_0)} x + \frac{(t - \Delta_0)}{(t + \Delta_0)}, \end{aligned} \quad (\text{A.349})$$

which is readily solved by

$$x_{\pm} = -\frac{\mu}{2(t + \Delta_0)} \pm \sqrt{\frac{\mu^2}{4(t + \Delta_0)^2} - \frac{(t - \Delta_0)}{(t + \Delta_0)}} = \frac{-\mu \pm \sqrt{\mu^2 - 4(t^2 - \Delta_0^2)}}{2(t + \Delta_0)}. \quad (\text{A.350})$$

These are precisely the decay parameters we encountered in Eq. (10.45). In order to motivate the solutions given in Eq. (10.44) we first note that Eq. (A.350) defines “quasi-momenta”

$$q_{\pm} := -i \ln(x_{\pm}) . \quad (\text{A.351})$$

with which

$$\gamma_{q_{\pm}}^{A/B} = \frac{1}{\sqrt{2L}} \sum_{j=1}^L e^{q_{\pm} R_j} \gamma_j^{A/B} = \frac{1}{\sqrt{2L}} \sum_{j=1}^L x_{\pm}^{R_j} \gamma_j^{A/B} \stackrel{(\diamond)}{=} \frac{1}{\sqrt{2L}} \sum_{j=1}^L x_{\pm}^j \gamma_j^{A/B} \quad (\text{A.352})$$

where we plugged in $R_j = a \cdot j \stackrel{a=1}{=} j$ at $(\diamond)$. Using this, make an ansatz

$$\Psi^{A/B}(q_{\pm}) = \psi_+^{A/B} \gamma_{q_+}^{A/B} + \psi_-^{A/B} \gamma_{q_-}^{A/B} , \quad (\text{A.353})$$

for $\Psi^{A/B}(\mathbf{k})$ from Eq. (A.342). Here, (ψ_+, ψ_-) are real coefficients. To arrive at a general solution, we will have to consider superpositions of the form

$$\begin{aligned} \Psi^A(q_{\pm}) &= \psi_+^A \gamma_{q_+}^A + \psi_-^A \gamma_{q_-}^A \\ &= \psi_+^A \frac{1}{\sqrt{2L}} \sum_{j=1}^L x_+^j \gamma_j^A + \psi_-^A \frac{1}{\sqrt{2L}} \sum_{j=1}^L x_-^j \gamma_j^A \\ &= \sum_{j=1}^L \left(\psi_+^A x_+^j + \psi_-^A x_-^j \right) \gamma_j^A \\ \Psi^B(q_{\pm}) &= \psi_+^B \gamma_{q_+}^B + \psi_-^B \gamma_{q_-}^B \\ &= \psi_+^B \frac{1}{\sqrt{2L}} \sum_{j=1}^L x_+^j \gamma_j^B + \psi_-^B \frac{1}{\sqrt{2L}} \sum_{j=1}^L x_-^j \gamma_j^B \\ &= \sum_{j=1}^L \left(\psi_+^B x_+^j + \psi_-^B x_-^j \right) \gamma_j^B , \end{aligned} \quad (\text{A.354})$$

where we have absorbed the prefactors $\frac{1}{\sqrt{2L}}$ into the real coefficients $\psi_{\pm}^{A/B}$ for better readability. Now, we are looking for two Majorana zero modes that are localised on opposite ends of the chain. Thus, we make the ansatz

$$\Gamma^A \equiv \sum_{j=1}^L \left(\alpha_+ x_+^j + \alpha_- x_-^j \right) \gamma_j^A \quad \text{and} \quad \Gamma^B \equiv \sum_{j=1}^L \left(\beta_+ x_+^j + \beta_- x_-^j \right) \gamma_{L+1-j}^B \quad (\text{A.355})$$

for our boundary Majorana zero modes Γ^A and Γ^B . These are normalisable for $L \rightarrow \infty$, and therefore viable solutions, if $|x_{\pm}| < 0$. In that case Γ^A is localised near $j = 1$ and Γ^B is localised near $j = L$. The coefficients (α_+, α_-) and (β_+, β_-) can be determined via the zero-mode constraint

$$[H_{\text{Kit}}, \Gamma^{A/B}] \stackrel{!}{=} 0 . \quad (\text{A.356})$$

Specifically, we find

$$\begin{aligned} [H_{\text{Kit}}, \Gamma^A] &= \sum_{j=1}^{L-1} \sum_{k=1}^L (\alpha_+ x_+^k + \alpha_- x_-^k) \left((\Delta_0 - t) [\gamma_j^A \gamma_{j+1}^B, \gamma_k^A] + (\Delta_0 + t) [\gamma_j^B \gamma_{j+1}^A, \gamma_k^A] \right) \\ &\quad - \mu \sum_{j,k=1}^L (\alpha_+ x_+^k + \alpha_- x_-^k) [\gamma_j^A \gamma_j^B, \gamma_k^A] \\ &\stackrel{(*)}{=} 2 \left[\sum_{j=1}^{L-1} \sum_{k=1}^L (\alpha_+ x_+^k + \alpha_- x_-^k) \left((t - \Delta_0) \delta_{jk} \gamma_{j+1}^B + (t + \Delta_0) \delta_{j+1,k} \gamma_j^B \right) + \mu \sum_{j,k=1}^L (\alpha_+ x_+^k + \alpha_- x_-^k) \delta_{jk} \gamma_j^B \right] \\ &\stackrel{(*)}{=} 2 \left[\sum_{j=1}^L \left\{ (\alpha_+ x_+^{j-1} + \alpha_- x_-^{j-1}) (t - \Delta_0) + (\alpha_+ x_+^{j+1} + \alpha_- x_-^{j+1}) (t + \Delta_0) + (\alpha_+ x_+^j + \alpha_- x_-^j) \mu \right\} \gamma_j^B \right. \\ &\quad \left. - (\alpha_+ x_+^0 + \alpha_- x_-^0) (t - \Delta_0) \gamma_1^B - (\alpha_+ x_+^{L+1} + \alpha_- x_-^{L+1}) (t + \Delta_0) \gamma_L^B \right] , \end{aligned} \quad (\text{A.357})$$

where we have used $\{\gamma_j^X, \gamma_k^Y\} = 2\delta_{jk}\delta_{XY}$ and $\{A, B\} = [A, B] - 2BA$ to rewrite

$$\begin{aligned}
[\gamma_j^X \gamma_k^Y, \gamma_l^Z] &= \gamma_j^X [\gamma_k^Y, \gamma_l^Z] + [\gamma_j^X, \gamma_l^Z] \gamma_k^Y \\
&= \gamma_j^X (\{\gamma_k^Y, \gamma_l^Z\} - 2\gamma_l^Z \gamma_k^Y) + (\{\gamma_j^X, \gamma_l^Z\} - 2\gamma_l^Z \gamma_j^X) \gamma_k^Y \\
&= \gamma_j^X \{\gamma_k^Y, \gamma_l^Z\} + \{\gamma_j^X, \gamma_l^Z\} \gamma_k^Y - 2(\gamma_l^Z \gamma_j^X + \gamma_j^X \gamma_l^Z) \gamma_k^Y \\
&= \gamma_j^X \{\gamma_k^Y, \gamma_l^Z\} + \{\gamma_j^X, \gamma_l^Z\} \gamma_k^Y - 2\{\gamma_j^X, \gamma_l^Z\} \gamma_k^Y \\
&= \gamma_j^X \{\gamma_k^Y, \gamma_l^Z\} - \{\gamma_j^X, \gamma_l^Z\} \gamma_k^Y \\
&= 2(\delta_{kl}\delta_{YZ}\gamma_j^X - \delta_{jl}\delta_{XZ}\gamma_k^Y)
\end{aligned} \tag{A.358}$$

such that

$$[\gamma_j^A \gamma_{j+1}^B, \gamma_k^A] = -2\delta_{jk}\gamma_{j+1}^B, \quad [\gamma_j^B \gamma_{j+1}^A, \gamma_k^A] = 2\delta_{j+1,k}\gamma_j^B, \quad [\gamma_j^A \gamma_j^B, \gamma_k^A] = -2\delta_{jk}\gamma_j^B \tag{A.359}$$

in $(\star)$. For $(*)$, we included additional terms $\propto \gamma_1^B$ and $\propto \gamma_L^B$ to get a simpler sum over all the sites $j = 1, \dots, L$ of the chain. For Eq. (A.357) to vanish, i.e. for Eq. (A.356) to be satisfied for Γ^A , we therefore require

$$\begin{aligned}
0 &\stackrel{!}{=} (\alpha_+ x_+^{j-1} + \alpha_- x_-^{j-1})(t - \Delta_0) + (\alpha_+ x_+^{j+1} + \alpha_- x_-^{j+1})(t + \Delta_0) + (\alpha_+ x_+^j + \alpha_- x_-^j)\mu \\
0 &\stackrel{!}{=} (\alpha_+ x_+^0 + \alpha_- x_-^0)(t - \Delta_0) \\
0 &\stackrel{!}{=} (\alpha_+ x_+^{L+1} + \alpha_- x_-^{L+1})(t + \Delta_0)
\end{aligned} \tag{A.360}$$

separately. If we rearrange the first constraint as

$$\begin{aligned}
0 &\stackrel{!}{=} (\alpha_+ x_+^{j-1} + \alpha_- x_-^{j-1})(t - \Delta_0) + (\alpha_+ x_+^{j+1} + \alpha_- x_-^{j+1})(t + \Delta_0) + (\alpha_+ x_+^j + \alpha_- x_-^j)\mu \\
&= \alpha_+ \left[(t + \Delta_0)x_+^{j+1} + \mu x_+^j + (t - \Delta_0)x_+^{j-1} \right] + \alpha_- \left[(t + \Delta_0)x_-^{j+1} + \mu x_-^j + (t - \Delta_0)x_-^{j-1} \right]
\end{aligned} \tag{A.361}$$

we find that it is automatically satisfied because

$$\begin{aligned}
&0 \stackrel{!}{=} (t + \Delta_0)x_{\pm}^{j+1} + \mu x_{\pm}^j + (t - \Delta_0)x_{\pm}^{j-1} \\
\iff &0 \stackrel{!}{=} (t + \Delta_0)x_{\pm}^2 + \mu x_{\pm} + (t - \Delta_0) \\
\iff &0 \stackrel{!}{=} x_{\pm}^2 + \frac{\mu}{(t + \Delta_0)}x_{\pm} + \frac{(t - \Delta_0)}{(t + \Delta_0)}
\end{aligned} \tag{A.362}$$

is readily solved by $x_{\pm}$ from Eq. (A.350). The coefficients can then be determined by the remaining constraints from Eq. (A.360). Concretely, since

$$0 \stackrel{!}{=} \alpha_+ x_+^{L+1} + \alpha_- x_-^{L+1} \tag{A.363}$$

is trivially satisfied for $L \rightarrow \infty$ when $|x_{\pm}| < 1$, the coefficients are determined by

$$0 \stackrel{!}{=} \alpha_+ x_+^0 + \alpha_- x_-^0. \tag{A.364}$$

An analogous calculation shows that the coefficients (β_+, β_-) of Γ^B are determined by

$$0 \stackrel{!}{=} \beta_+ x_+^0 + \beta_- x_-^0. \tag{A.365}$$

Torus Gauge-Equivalence to the X-Gate

Here we demonstrate the validity of Eq. (10.124) through explicit calculation:

$$\begin{aligned}
U\mathcal{U}_{\text{WZ}}(C^{0,1})U^\dagger e^{i\delta} &= \begin{pmatrix} e^{i\varphi_1} & 0 & 0 & 0 \\ 0 & e^{i\varphi_2} & 0 & 0 \\ 0 & 0 & e^{i\varphi_3} & 0 \\ 0 & 0 & 0 & e^{i\varphi_4} \end{pmatrix} e^{i\theta} \begin{pmatrix} 0 & e^{-i\alpha} & 0 & 0 \\ e^{i\alpha} & 0 & 0 & 0 \\ 0 & 0 & 0 & e^{-i\beta} \\ 0 & 0 & e^{i\beta} & 0 \end{pmatrix} \begin{pmatrix} e^{-i\varphi_1} & 0 & 0 & 0 \\ 0 & e^{-i\varphi_2} & 0 & 0 \\ 0 & 0 & e^{-i\varphi_3} & 0 \\ 0 & 0 & 0 & e^{-i\varphi_4} \end{pmatrix} e^{i\delta} \\
&\stackrel{(\diamond)}{=} \begin{pmatrix} e^{i\varphi_1} & 0 & 0 & 0 \\ 0 & e^{-i\varphi_1} & 0 & 0 \\ 0 & 0 & e^{i\varphi_3} & 0 \\ 0 & 0 & 0 & e^{-i\varphi_3} \end{pmatrix} e^{i\theta} \begin{pmatrix} 0 & e^{-i\alpha} & 0 & 0 \\ e^{i\alpha} & 0 & 0 & 0 \\ 0 & 0 & 0 & e^{-i\beta} \\ 0 & 0 & e^{i\beta} & 0 \end{pmatrix} \begin{pmatrix} e^{-i\varphi_1} & 0 & 0 & 0 \\ 0 & e^{i\varphi_1} & 0 & 0 \\ 0 & 0 & e^{-i\varphi_3} & 0 \\ 0 & 0 & 0 & e^{i\varphi_3} \end{pmatrix} e^{i\delta} \\
&\stackrel{(*)}{=} \begin{pmatrix} 0 & e^{-i(\alpha-2\varphi_1)} & 0 & 0 \\ e^{i(\alpha-2\varphi_1)} & 0 & 0 & 0 \\ 0 & 0 & 0 & e^{-i(\beta-2\varphi_3)} \\ 0 & 0 & e^{i(\beta-2\varphi_3)} & 0 \end{pmatrix} e^{i(\delta+\theta)} \\
&\stackrel{(*)}{=} \begin{pmatrix} 0 & e^{-i(\alpha-\alpha)} & 0 & 0 \\ e^{i(\alpha-\alpha)} & 0 & 0 & 0 \\ 0 & 0 & 0 & e^{-i(\beta-\beta+\pi)} \\ 0 & 0 & e^{i(\beta-\beta+\pi)} & 0 \end{pmatrix} e^{i\frac{\pi}{2}} \\
&= \begin{pmatrix} 0 & i & 0 & 0 \\ i & 0 & 0 & 0 \\ 0 & 0 & 0 & -i \\ 0 & 0 & -i & 0 \end{pmatrix}, \tag{A.366}
\end{aligned}$$

where we used $\varphi_1 = -\varphi_2$ and $\varphi_3 = -\varphi_4$ in $(\diamond)$ and plugged in $\delta = \pi/2 - \theta$, $\varphi_1 = \alpha/2$ and $\varphi_3 = (\beta - \pi)/2$ in $(*)$. If we instead plug in $\delta = -\theta$, $\varphi_1 = -\varphi_2 = \alpha/2$ and $\varphi_3 = -\varphi_4 = \beta/2$ in $(*)$ we readily get

$$U\mathcal{U}_{\text{WZ}}(C^{0,1})U^\dagger e^{i\delta} = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}. \tag{A.367}$$

B – Bibliography

- [1] E. P. Wigner, The unreasonable effectiveness of mathematics in the natural sciences, *Commun. Pure Appl. Math* **13**, 1 (1960).
- [2] P. A. M. Dirac, Quantised singularities in the electromagnetic field, *Proc. R. Soc. London A* **133**, 60 (1931).
- [3] Y. Aharonov and D. Bohm, Significance of electromagnetic potentials in the quantum theory, *Phys. Rev.* **115**, 485 (1959).
- [4] T. H. R. Skyrme and B. F. J. Schonland, A non-linear field theory, *Proc. R. Soc. Lond. A* **260**, 127 (1961).
- [5] T. Skyrme, A unified field theory of mesons and baryons, *Nucl. Phys.* **31**, 556 (1962).
- [6] G. 't Hooft, Magnetic monopoles in unified gauge theories, *Nucl. Phys. B* **79**, 276 (1974).
- [7] A. Belavin, A. Polyakov, A. Schwartz, and Y. Tyupkin, Pseudoparticle solutions of the Yang–Mills equations, *Phys. Rev. B* **59**, 85 (1975).
- [8] G. 't Hooft, Symmetry breaking through Bell–Jackiw anomalies, *Phys. Rev. Lett.* **37**, 8 (1976).
- [9] K. v. Klitzing, G. Dorda, and M. Pepper, New method for high-accuracy determination of the fine-structure constant based on quantized Hall resistance, *Phys. Rev. Lett.* **45**, 494 (1980).
- [10] D. J. Thouless, M. Kohmoto, M. P. Nightingale, and M. den Nijs, Quantized Hall conductance in a two-dimensional periodic potential, *Phys. Rev. Lett.* **49**, 405 (1982).
- [11] D. C. Tsui, H. L. Stormer, and A. C. Gossard, Two-dimensional magnetotransport in the extreme quantum limit, *Phys. Rev. Lett.* **48**, 1559 (1982).
- [12] M. F. Atiyah, Topological quantum field theories, *Publ. Math. IHÉS* **68**, 175 (1988).
- [13] E. Witten, Topological quantum field theory, *Commun. Math. Phys.* **117**, 353 (1988).
- [14] C.-K. Chiu, J. C. Y. Teo, A. P. Schnyder, and S. Ryu, Classification of topological quantum matter with symmetries, *Rev. Mod. Phys.* **88**, 035005 (2016).
- [15] A. Hatcher, *Vector Bundles and K-Theory* (2003), available at URL.
- [16] J. Bellissard, A. van Elst, and H. Schulz-Baldes, The noncommutative geometry of the quantum Hall effect, *J. Math. Phys.* **35**, 5373 (1994).
- [17] E. Prodan and H. Schulz-Baldes, *Bulk and Boundary Invariants for Complex Topological Insulators: From K-Theory to Physics*, Mathematical Physics Studies (Springer, Cham, 2016).
- [18] E. Prodan, The non-commutative geometry of the complex classes of topological insulators, *Topol. Quantum Matter* **1** (2014).
- [19] F. D. M. Haldane, Model for a quantum Hall effect without Landau levels: condensed-matter realization of the “parity anomaly”, *Phys. Rev. Lett.* **61**, 2015 (1988).
- [20] C. L. Kane and E. J. Mele, Quantum spin Hall effect in graphene, *Phys. Rev. Lett.* **95**, 226801 (2005).

-
- [21] C. L. Kane and E. J. Mele, Z_2 topological order and the quantum spin Hall effect, *Phys. Rev. Lett.* **95**, 146802 (2005).
 - [22] X. Wan, A. M. Turner, A. Vishwanath, and S. Y. Savrasov, Topological semimetal and Fermi-arc surface states in the electronic structure of pyrochlore iridates, *Phys. Rev. B* **83**, 205101 (2011).
 - [23] A. A. Burkov and L. Balents, Weyl semimetal in a topological insulator multilayer, *Phys. Rev. Lett.* **107**, 127205 (2011).
 - [24] H. Nielsen and M. Ninomiya, A no-go theorem for regularizing chiral fermions, *Phys. Lett. B* **105**, 219 (1981).
 - [25] H. Fukaya, M. Furuta, S. Matsuo, T. Onogi, S. Yamaguchi, and M. Yamashita, A physicist-friendly reformulation of the Atiyah–Patodi–Singer index and its mathematical justification, arXiv:2001.01428.
 - [26] M. R. Zirnbauer, Riemannian symmetric superspaces and their origin in random-matrix theory, *J. Math. Phys.* **37**, 4986 (1996).
 - [27] A. Altland and M. R. Zirnbauer, Nonstandard symmetry classes in mesoscopic normal-superconducting hybrid structures, *Phys. Rev. B* **55**, 1142 (1997).
 - [28] J. S. Russell, in *Report of the Fourteenth Meeting of the British Association for the Advancement of Science*, pp. 311–390 (British Association for the Advancement of Science, 1844), available at URL.
 - [29] N. J. Zabusky and M. D. Kruskal, Interaction of “solitons” in a collisionless plasma and the recurrence of initial states, *Phys. Rev. Lett.* **15**, 240 (1965).
 - [30] N. Manton and P. Sutcliffe, *Topological Solitons*, (Cambridge University Press, Cambridge, United Kingdom, 2004).
 - [31] A. I. Vainshtein, V. I. Zakharov, V. A. Novikov, and M. A. Shifman, Abc of instantons, *Sov. Phys. Usp.* **25**, 195 (1982).
 - [32] M. A. Shifman, ed., *Instantons in Gauge Theories*, (World Scientific, Singapore, 1994).
 - [33] N. S. Kiselev, A. N. Bogdanov, R. Schäfer, and U. K. Rößler, Chiral skyrmions in thin magnetic films: new objects for magnetic storage technologies?, *J. Phys. D: Appl. Phys.* **44**, 392001 (2011).
 - [34] J. S. Alden, A. W. Tsen, P. Y. Huang, R. Hovden, L. Brown, J. Park, D. A. Muller, and P. L. McEuen, Strain solitons and topological defects in bilayer graphene, *Proc. Natl. Acad. Sci. U.S.A* **110**, 11256–11260 (2013).
 - [35] S. Zhang, Q. Xu, Y. Hou, A. Song, Y. Ma, L. Gao, M. Zhu, T. Ma, L. Liu, X.-Q. Feng, and Q. Li, Domino-like stacking order switching in twisted monolayer–multilayer graphene, *Nat. Mater.* **21**, 621 (2022).
 - [36] L. Onsager, Statistical hydrodynamics, *Il Nuovo Cimento* **6**, 279 (1949).
 - [37] A. A. Abrikosov, On the magnetic properties of superconductors of the second group, *Sov. Phys. JETP* **5**, 1174 (1957), available at URL.
 - [38] R. Auzzi, S. Bolognesi, J. Evslin, K. Konishi, and A. Yung, Nonabelian superconductors: vortices and confinement in $N = 2$ SQCD, *Nucl. Phys. B* **673**, 187–216 (2003).
 - [39] M. Nakahara, *Geometry, Topology and Physics*, (Institute of Physics Publishing, Bristol, United Kingdom, 2003).
 - [40] A. M. Polyakov, Particle spectrum in quantum field theory, *JETP Lett.* **20**, 194 (1974), available at URL.

-
- [41] J. M. D. Coey, *Magnetism and Magnetic Materials*, (Cambridge University Press, Cambridge, United Kingdom, 2010).
 - [42] A. M. Kosevich, V. V. Gann, A. I. Zhukov, and V. P. Voronov, Magnetic soliton motion in a nonuniform magnetic field, *Sov. Phys. JETP* **87**, 401 (1998).
 - [43] J. R. Taylor, *Optical Solitons: Theory and Experiment*, (Cambridge University Press, Cambridge, United Kingdom, 1992).
 - [44] M. R. R. Vaziri, Describing the propagation of intense laser pulses in nonlinear Kerr media using the ducting model, *Laser Phys.* **23**, 105401 (2013).
 - [45] J. W. Miles, Solitary waves, *Annu. Rev. Fluid Mech.* **12**, 11 (1980).
 - [46] C. von Westenholz, Topological and Noether-conservation laws, *Ann. Inst. H. Poincaré Phys. Théor.* **30**, 353 (1979), available at URL.
 - [47] D. S. Freed, The Atiyah–Singer index theorem, arXiv:2107.03557.
 - [48] M. F. Atiyah, V. K. Patodi, and I. M. Singer, Spectral asymmetry and Riemannian geometry. i, *Math. Proc. Camb. Philos. Soc.* **77**, 43–69 (1975).
 - [49] M. F. Atiyah, V. K. Patodi, and I. M. Singer, Spectral asymmetry and Riemannian geometry. ii, *Math. Proc. Camb. Philos. Soc.* **78**, 405–432 (1975).
 - [50] M. F. Atiyah, V. K. Patodi, and I. M. Singer, Spectral asymmetry and Riemannian geometry. iii, *Math. Proc. Camb. Philos. Soc.* **79**, 71–99 (1976).
 - [51] P. Delplace, J. B. Marston, and A. Venaille, Topological origin of equatorial waves, *Science* **358**, 1075 (2017).
 - [52] D. Tong, A gauge theory for shallow water, *SciPost Phys.* **14**, 102 (2023).
 - [53] K. Fujikawa, Path-integral measure for gauge-invariant fermion theories, *Phys. Rev. Lett.* **42**, 1195 (1979).
 - [54] K. Fujikawa, Path integral for gauge theories with fermions, *Phys. Rev. D* **21**, 2848 (1980).
 - [55] C. Callan and J. Harvey, Anomalies and fermion zero modes on strings and domain walls, *Nucl. Phys. B* **250**, 427 (1985).
 - [56] L. Alvarez-Gaumé and E. Witten, Gravitational anomalies, *Nucl. Phys. B* **234**, 269 (1984).
 - [57] S. L. Adler, Axial-vector vertex in spinor electrodynamics, *Phys. Rev.* **177**, 2426 (1969).
 - [58] J. S. Bell and R. Jackiw, A pcac puzzle: $\pi^0 \rightarrow \gamma\gamma$ in the σ model, *Il Nuovo Cimento A* **60**, 47 (1969).
 - [59] C. G. Callan, R. Dashen, and D. J. Gross, Toward a theory of the strong interactions, *Phys. Rev. D* **17**, 2717 (1978).
 - [60] G. Galilei, *The Assayer (Il Saggiatore)* (1623), available at URL.
 - [61] A. Hatcher, *Algebraic Topology* (Cambridge University Press, Cambridge, United Kingdom, 2002), available at URL.
 - [62] B. Riemann, Theorie der Abel’schen Funktionen, *Journal für die reine und angewandte Mathematik* **54**, 101 (1857), available at URL.
 - [63] Steelpillow, Sphercycles1 (2016), published via Wikimedia Commons, licensed under Creative Commons Attribution-ShareAlike 4.0 International (CC BY-SA 4.0).

-
- [64] Steelpillow, Toruscycles1 (2016), published via Wikimedia Commons, licensed under Creative Commons Attribution-ShareAlike 4.0 International (CC BY-SA 4.0).
- [65] M. Actipes, On the fundamental group of surfaces (2013), available at URL.
- [66] H. Toda, *Composition Methods in Homotopy Groups of Spheres*, (Princeton University Press, Princeton, NJ, 1962).
- [67] R. Abraham and J. E. Marsden, *Foundations of Mechanics* (Addison-Wesley, Reading, MA, 1987), available at URL.
- [68] M. Fruchart, *Topological Phases of Periodically Driven Crystals*, Ph.D. thesis, Université de Lyon (2016), available at URL.
- [69] H. Whitney, The self-intersections of a smooth n -manifold in $2n$ -space, *Ann. Math.* **45**, 220 (1944).
- [70] D. Bleecker, *Gauge Theory and Variational Principles* (Dover Publications, Mineola, NY, USA, 2005), available at URL.
- [71] Tazerenix, Ehresmann connection (2021), published via Wikimedia Commons, licensed under Creative Commons Attribution-ShareAlike 4.0 International (CC BY-SA 4.0).
- [72] V. Pati, Obstruction theory and characteristic classes (2006), available at URL.
- [73] Nobel Prize Outreach, The nobel prize in physics 2016 (2016), available at URL.
- [74] M. F. Atiyah and I. M. Singer, The index of elliptic operators I, *Ann. Math.* **87**, 484 (1968).
- [75] H. Nielsen and M. Ninomiya, The Adler–Bell–Jackiw anomaly and Weyl fermions in a crystal, *Phys. Lett. B* **130**, 389 (1983).
- [76] E. Wigner, *Gruppentheorie und ihre Anwendung auf die Quantenmechanik der Atomspektren*, (Vieweg+Teubner Verlag Wiesbaden, Braunschweig, Germany, 1931).
- [77] J.-P. Serre, *Linear Representations of Finite Groups*, Graduate Texts in Mathematics (Springer, Berlin, 1977).
- [78] M. R. Zirnbauer, Particle–hole symmetries in condensed matter, *J. Math. Phys.* **62** (2021).
- [79] Y. Aharonov and J. Anandan, Phase change during a cyclic quantum evolution, *Phys. Rev. Lett.* **58**, 1593 (1987).
- [80] A. Bohm, A. Mostafazadeh, H. Koizumi, Q. Niu, and J. Zwanziger, *The Geometric Phase in Quantum Systems*, Theoretical and Mathematical Physics (Springer, Berlin, Heidelberg, 2003).
- [81] M. V. Berry, Quantal phase factors accompanying adiabatic changes, *Proc. R. Soc. London A* **392**, 45 (1984).
- [82] C. Timm, Theory of superconductivity (2020), TU Dresden, available at URL.
- [83] P. Ring and P. Schuck, *The Nuclear Many-Body Problem* (Springer Berlin Heidelberg, Berlin, Heidelberg, 2004), available at URL.
- [84] J. Bardeen, L. N. Cooper, and J. R. Schrieffer, Theory of superconductivity, *Physical Review* **108**, 1175 (1957).
- [85] R. M. Fernandes, Lecture notes: BCS theory of superconductivity (2012), available at URL.
- [86] C. Kallin and J. Berlinsky, Chiral superconductors, *Rep. Prog. Phys.* **79**, 054502 (2016).
- [87] K. Neergård, Fock space representations of Bogolyubov transformations as spin representations, *Phys. Rev. C* **107**, 024315 (2023).

-
- [88] C. Bloch and A. Messiah, The canonical form of an antisymmetric tensor and its application to the theory of superconductivity, *Nucl. Phys.* **39**, 95 (1962).
 - [89] N. Onishi and S. Yoshida, Generator coordinate method applied to nuclei in the transition region, *Nucl. Phys.* **80**, 367 (1966).
 - [90] L. M. Robledo, Technical aspects of the evaluation of the overlap of Hartree–Fock–Bogoliubov wave functions, *Phys. Rev. C* **84**, 014307 (2011).
 - [91] L. M. Robledo, Practical formulation of the extended Wick’s theorem and the Onishi formula, *Phys. Rev. C* **50**, 2874 (1994).
 - [92] L. M. Robledo, Sign of the overlap of Hartree–Fock–Bogoliubov wave functions, *Phys. Rev. C* **79**, 021302 (2009).
 - [93] G. F. Bertsch and L. M. Robledo, Symmetry restoration in Hartree–Fock–Bogoliubov based theories, *Phys. Rev. Lett.* **108**, 042505 (2012).
 - [94] E. Mascot, T. Hodge, D. Crawford, J. Bedow, D. K. Morr, and S. Rachel, Many-body Majorana braiding without an exponential Hilbert space, *Phys. Rev. Lett.* **131**, 176601 (2023).
 - [95] S. Rao, *Introduction to Abelian and Non-Abelian Anyons*, pp. 399–437, (Springer Singapore, Singapore, 2017).
 - [96] E. Artin, Theorie der Zöpfe, *Abhandlungen aus dem Mathematischen Seminar der Universität Hamburg* **4**, 47 (1925).
 - [97] E. Artin, Theory of braids, *Ann. Math.* **48**, 101 (1947).
 - [98] F. Wilczek, Quantum mechanics of fractional-spin particles, *Phys. Rev. Lett.* **49**, 957 (1982).
 - [99] G. A. Goldin and D. H. Sharp, Rotation generators in two-dimensional space and particles obeying unusual statistics, *Phys. Rev. D* **28**, 830 (1983).
 - [100] E. C. Rowell and Z. Wang, Mathematics of topological quantum computing, *Bull. Am. Math. Soc.* **55**, 183 (2018).
 - [101] C. Nayak, S. H. Simon, A. Stern, M. Freedman, and S. Das Sarma, Non-Abelian anyons and topological quantum computation, *Rev. Mod. Phys.* **80**, 1083 (2008).
 - [102] V. Lahtinen and J. K. Pachos, A Short Introduction to Topological Quantum Computation, *SciPost Phys.* **3**, 021 (2017).
 - [103] S. Trebst, M. Troyer, Z. Wang, and A. W. W. Ludwig, A short introduction to fibonacci anyon models, *Prog. Theor. Phys. Supp.* **176**, 384 (2008).
 - [104] E. Rowell, R. Stong, and Z. Wang, On classification of modular tensor categories, *Commun. Math. Phys.* **292**, 343 (2009).
 - [105] S. Autti, V. V. Dmitriev, J. T. Mäkinen, A. A. Soldatov, G. E. Volovik, A. N. Yudin, V. V. Zavjalov, and V. B. Eltsov, Observation of half-quantum vortices in topological superfluid ^3He , *Phys. Rev. Lett.* **117**, 255301 (2016).
 - [106] D. A. Ivanov, Non-Abelian statistics of half-quantum vortices in p -wave superconductors, *Phys. Rev. Lett.* **86**, 268 (2001).
 - [107] G. E. Volovik, Fermion zero modes on vortices in chiral superconductors, *JETP Lett.* **70**, 609–614 (1999).
 - [108] J. Alicea, Y. Oreg, G. Refael, F. von Oppen, and M. P. A. Fisher, Non-Abelian statistics and topological quantum information processing in 1D wire networks, *Nat. Phys.* **7**, 412 (2011).

-
- [109] T. H. Hansson, A. Karlhede, and M. Sato, Topological field theory for p -wave superconductors, *New J. Phys.* **14**, 063017 (2012).
 - [110] F. Wilczek, Projective statistics and spinors in Hilbert space, arXiv:hep-th/9806228.
 - [111] M. Freedman, M. B. Hastings, C. Nayak, X.-L. Qi, K. Walker, and Z. Wang, Projective ribbon permutation statistics: a remnant of non-Abelian braiding in higher dimensions, *Phys. Rev. B* **83**, 115132 (2011).
 - [112] Y.-Z. You and X.-G. Wen, Projective non-Abelian statistics of dislocation defects in a $\mathbb{Z}_N$ rotor model, *Phys. Rev. B* **86**, 161107 (2012).
 - [113] N. Friis, Reasonable fermionic quantum information theories require relativity, *New J. Phys.* **18**, 033014 (2016).
 - [114] V. Lahtinen and J. K. Pachos, A Short Introduction to Topological Quantum Computation, *SciPost Phys.* **3**, 021 (2017).
 - [115] J. R. Wootton, J. Burri, S. Iblisdir, and D. Loss, Error correction for non-Abelian topological quantum computation, *Phys. Rev. X* **4**, 011051 (2014).
 - [116] V. Lahtinen and J. K. Pachos, Non-Abelian statistics as a Berry phase in exactly solvable models, *New J. Phys.* **11**, 093027 (2009).
 - [117] M. Cheng, V. Galitski, and S. Das Sarma, Nonadiabatic effects in the braiding of non-Abelian anyons in topological superconductors, *Phys. Rev. B* **84**, 104529 (2011).
 - [118] A. Feiguin, S. Trebst, A. W. W. Ludwig, M. Troyer, A. Kitaev, Z. Wang, and M. H. Freedman, Interacting anyons in topological quantum liquids: The golden chain, *Phys. Rev. Lett.* **98**, 160409 (2007).
 - [119] C. Gils, E. Ardonne, S. Trebst, A. W. W. Ludwig, M. Troyer, and Z. Wang, Collective states of interacting anyons, edge states, and the nucleation of topological liquids, *Phys. Rev. Lett.* **103**, 070401 (2009).
 - [120] A. W. W. Ludwig, D. Poilblanc, S. Trebst, and M. Troyer, Two-dimensional quantum liquids from interacting non-Abelian anyons, *New J. Phys.* **13**, 045014 (2011).
 - [121] V. Lahtinen, A. W. W. Ludwig, J. K. Pachos, and S. Trebst, Topological liquid nucleation induced by vortex-vortex interactions in Kitaev's honeycomb model, *Phys. Rev. B* **86**, 075115 (2012).
 - [122] G. Kells, V. Lahtinen, and J. Vala, Kitaev spin models from topological nanowire networks, *Phys. Rev. B* **89**, 075122 (2014).
 - [123] D. Wang, Z. Huang, and C. Wu, Fate and remnants of Majorana zero modes in a quantum wire array, *Phys. Rev. B* **89**, 174510 (2014).
 - [124] T. Liu and M. Franz, Electronic structure of topological superconductors in the presence of a vortex lattice, *Phys. Rev. B* **92**, 134519 (2015).
 - [125] J. A. Sobota, Y. He, and Z.-X. Shen, Angle-resolved photoemission studies of quantum materials, *Rev. Mod. Phys.* **93**, 025006 (2021).
 - [126] J. W. González and J. Fernández-Rossier, Impurity states in the quantum spin-Hall phase in graphene, *Phys. Rev. B* **86**, 115327 (2012).
 - [127] A. M. Black-Schaffer and A. V. Balatsky, Subsurface impurities and vacancies in a three-dimensional topological insulator, *Phys. Rev. B* **86**, 115433 (2012).
 - [128] C. Chan and T.-K. Ng, Impurity scattering in the bulk of topological insulators, *Phys. Rev. B* **85**, 115207 (2012).

-
- [129] W.-Y. Shan, J. Lu, H.-Z. Lu, and S.-Q. Shen, Vacancy-induced bound states in topological insulators, *Phys. Rev. B* **84**, 035307 (2011).
 - [130] J. Lu, W.-Y. Shan, H.-Z. Lu, and S.-Q. Shen, Non-magnetic impurities and in-gap bound states in topological insulators, *New J. Phys.* **13**, 103016 (2011).
 - [131] R.-J. Slager, L. Rademaker, J. Zaanen, and L. Balents, Impurity-bound states and Green's function zeros as local signatures of topology, *Phys. Rev. B* **92**, 085126 (2015).
 - [132] V. B. Jha, G. Rani, and R. Ganesh, Impurity-induced current in a Chern insulator, *Phys. Rev. B* **95**, 115434 (2017).
 - [133] L. Kimme and T. Hyart, Existence of zero-energy impurity states in different classes of topological insulators and superconductors and their relation to topological phase transitions, *Phys. Rev. B* **93**, 035134 (2016).
 - [134] Y. Liu and S. Chen, Diagnosis of bulk phase diagram of nonreciprocal topological lattices by impurity modes, *Phys. Rev. B* **102**, 075404 (2020).
 - [135] S.-S. Diop, L. Fritz, M. Vojta, and S. Rachel, Impurity bound states as detectors of topological band structures revisited, *Phys. Rev. B* **101**, 245132 (2020).
 - [136] S. Michel and M. Potthoff, Spin Berry curvature of the Haldane model, *Phys. Rev. B* **106**, 235423 (2022).
 - [137] B. Simon, Holonomy, the quantum adiabatic theorem, and Berry's phase, *Phys. Rev. Lett.* **51**, 2167 (1983).
 - [138] E. A. Fajardo and R. Winkler, Effective dynamics of two-dimensional Bloch electrons in external fields derived from symmetry, *Phys. Rev. B* **100**, 125301 (2019).
 - [139] J. C. Y. Teo and C. L. Kane, Topological defects and gapless modes in insulators and superconductors, *Phys. Rev. B* **82**, 115120 (2010).
 - [140] B. E. Kane, A silicon-based nuclear spin quantum computer, *Nature* **398**, 133 (1998).
 - [141] V. Lahtinen and J. K. Pachos, A short introduction to topological quantum computation, *SciPost Phys.* **3**, 021 (2017).
 - [142] S. Das Sarma, M. Freedman, and C. Nayak, Majorana zero modes and topological quantum computation, *npj Quantum Inf.* **1**, 15001 (2015).
 - [143] S. A. Wolf, D. D. Awschalom, R. A. Buhrman, J. M. Daughton, S. von Molnár, M. L. Roukes, A. Y. Chtchelkanova, and D. M. Treger, Spintronics: A spin-based electronics vision for the future, *Science* **294**, 1488 (2001).
 - [144] A. R. Akhmerov, C. W. Groth, J. Tworzydło, and C. W. J. Beenakker, Switching of electrical current by spin precession in the first Landau level of an inverted-gap semiconductor, *Phys. Rev. B* **80**, 195320 (2009).
 - [145] T. Yokoyama, Y. Tanaka, and N. Nagaosa, Giant spin rotation in the junction between a normal metal and a quantum spin Hall system, *Phys. Rev. Lett.* **102**, 166801 (2009).
 - [146] J. S. Van Dyke and D. K. Morr, Controlling the flow of spin and charge in nanoscopic topological insulators, *Phys. Rev. B* **93**, 081401 (2016).
 - [147] A. Narayan, A. Hurley, and S. Sanvito, Gate controlled spin pumping at a quantum spin Hall edge, *Appl. Phys. Lett.* **103**, 142407 (2013).
 - [148] A. Hurley, A. Narayan, and S. Sanvito, Spin-pumping and inelastic electron tunneling spectroscopy in topological insulators, *Phys. Rev. B* **87**, 245410 (2013).

-
- [149] C.-Z. Chang, P. Wei, and J. S. Moodera, Breaking time reversal symmetry in topological insulators, *MRS Bull.* **39**, 867D872 (2014).
- [150] J. Maciejko, C. Liu, Y. Oreg, X.-L. Qi, C. Wu, and S.-C. Zhang, Kondo effect in the helical edge liquid of the quantum spin Hall state, *Phys. Rev. Lett.* **102**, 256803 (2009).
- [151] Y. Tanaka, A. Furusaki, and K. A. Matveev, Conductance of a helical edge liquid coupled to a magnetic impurity, *Phys. Rev. Lett.* **106**, 236402 (2011).
- [152] B. A. Bernevig, T. L. Hughes, and S.-C. Zhang, Quantum spin Hall effect and topological phase transition in HgTe quantum wells, *Science* **314**, 1757 (2006).
- [153] M. König, S. Wiedmann, C. Brüne, A. Roth, H. Buhmann, L. W. Molenkamp, X.-L. Qi, and S.-C. Zhang, Quantum spin Hall insulator state in HgTe quantum wells, *Science* **318**, 766 (2007).
- [154] C.-C. Liu, H. Jiang, and Y. Yao, Low-energy effective Hamiltonian involving spin-orbit coupling in silicene and two-dimensional germanium and tin, *Phys. Rev. B* **84**, 195430 (2011).
- [155] M. Ezawa, Monolayer topological insulators: Silicene, germanene, and stanene, *J. Phys. Soc. Jpn.* **84**, 121003 (2015).
- [156] J. M. Pizarro, S. Adler, K. Zantout, T. Mertz, P. Barone, R. Valentí, G. Sangiovanni, and T. O. Wehling, Deconfinement of Mott localized electrons into topological and spin-orbit-coupled Dirac fermions, *npj Quantum Mater.* **5** (2020).
- [157] R. Yu, W. Zhang, H.-J. Zhang, S.-C. Zhang, X. Dai, and Z. Fang, Quantized anomalous Hall effect in magnetic topological insulators, *Science* **329**, 61 (2010).
- [158] Y. L. Chen, J.-H. Chu, J. G. Analytis, Z. K. Liu, K. Igarashi, H.-H. Kuo, X. L. Qi, S. K. Mo, R. G. Moore, D. H. Lu, M. Hashimoto, T. Sasagawa, S. C. Zhang, I. R. Fisher, Z. Hussain, and Z. X. Shen, Massive Dirac fermion on the surface of a magnetically doped topological insulator, *Science* **329**, 659 (2010).
- [159] T. Eelbo, M. Waśniowska, M. Sikora, M. Dobrzański, A. Kozłowski, A. Pulkin, G. Autès, I. Miotkowski, O. V. Yazyev, and R. Wiesendanger, Strong out-of-plane magnetic anisotropy of Fe adatoms on Bi₂Te₃, *Phys. Rev. B* **89**, 104424 (2014).
- [160] M. Vondráček, L. Cornils, J. Minár, J. Warmuth, M. Michiardi, C. Piamonteze, L. Barreto, J. A. Miwa, M. Bianchi, P. Hofmann, L. Zhou, A. Kamlapure, A. A. Khajetoorians, R. Wiesendanger, J.-L. Mi, B.-B. Iversen, S. Mankovsky, S. Borek, H. Ebert, M. Schüler, T. Wehling, J. Wiebe, and J. Honolka, Nickel: The time-reversal symmetry conserving partner of iron on a chalcogenide topological insulator, *Phys. Rev. B* **94**, 161114 (2016).
- [161] X.-L. Qi, R. Li, J. Zang, and S.-C. Zhang, Inducing a magnetic monopole with topological surface states, *Science* **323**, 1184 (2009).
- [162] C. Kittel and P. McEuen, *Introduction to Solid State Physics* (Wiley, New Delhi, 2015), available at URL.
- [163] Q. Liu, C.-X. Liu, C. Xu, X.-L. Qi, and S.-C. Zhang, Magnetic impurities on the surface of a topological insulator, *Phys. Rev. Lett.* **102**, 156603 (2009).
- [164] J. Gao, W. Chen, X. C. Xie, and F.-c. Zhang, In-plane noncollinear exchange coupling mediated by helical edge states in quantum spin Hall systems, *Phys. Rev. B* **80**, 241302 (2009).
- [165] N. Bloembergen and T. J. Rowland, Nuclear spin exchange in solids: Tl²⁰³ and Tl²⁰⁵ magnetic resonance in thallium and thallic oxide, *Phys. Rev.* **97**, 1679 (1955).
- [166] V. D. Kurilovich, P. D. Kurilovich, and I. S. Burmistrov, Indirect exchange interaction between magnetic impurities near the helical edge, *Phys. Rev. B* **95**, 115430 (2017).

-
- [167] M. V. Hosseini, Z. Karimi, and J. Davoodi, Indirect exchange interaction between magnetic impurities in one-dimensional gapped helical states, *J. Phys.: Condens. Matter* **33**, 085801 (2020).
 - [168] O. M. Yevtushenko and V. I. Yudson, Kondo impurities coupled to a helical Luttinger liquid: RKKY-Kondo physics revisited, *Phys. Rev. Lett.* **120**, 147201 (2018).
 - [169] J. Bouaziz, M. d. S. Dias, F. S. M. Guimarães, and S. Lounis, Spin dynamics of $3d$ and $4d$ impurities embedded in prototypical topological insulators, *Phys. Rev. Mater.* **3**, 054201 (2019).
 - [170] S. Pradhan and J. Fransson, Time-dependent potential impurity in a topological insulator, *Phys. Rev. B* **100**, 125163 (2019).
 - [171] T. Tatsumi, H. Yoshioka, and M. Hayashi, Transport properties of zigzag ribbons composed of two-dimensional quantum spin Hall insulators with a single magnetic impurity, *J. Phys. Soc. Jpn.* **89**, 043702 (2020).
 - [172] M. Elbracht and M. Potthoff, Long-time relaxation dynamics of a spin coupled to a Chern insulator, *Phys. Rev. B* **103**, 024301 (2021).
 - [173] U. Bajpai and B. K. Nikolić, Time-retarded damping and magnetic inertia in the Landau-Lifshitz-Gilbert equation self-consistently coupled to electronic time-dependent nonequilibrium Green functions, *Phys. Rev. B* **99**, 134409 (2019).
 - [174] M. Onoda and N. Nagaosa, Dynamics of localized spins coupled to the conduction electrons with charge and spin currents, *Phys. Rev. Lett.* **96**, 066603 (2006).
 - [175] S. Bhattacharjee, L. Nordström, and J. Fransson, Atomistic spin dynamic method with both damping and moment of inertia effects included from first principles, *Phys. Rev. Lett.* **108**, 057204 (2012).
 - [176] R. Wiesendanger, Spin mapping at the nanoscale and atomic scale, *Rev. Mod. Phys.* **81**, 1495 (2009).
 - [177] L. Zhou, J. Wiebe, S. Lounis, E. Vedmedenko, F. Meier, S. Blügel, P. H. Dederichs, and R. Wiesendanger, Strength and directionality of surface Ruderman-Kittel-Kasuya-Yosida interaction mapped on the atomic scale, *Nat. Physics* **6**, 187 (2010).
 - [178] A. van Houselt and H. J. W. Zandvliet, Colloquium: Time-resolved scanning tunneling microscopy, *Rev. Mod. Phys.* **82**, 1593 (2010).
 - [179] Y. Tian, F. Yang, C. Guo, and Y. Jiang, Recent advances in ultrafast time-resolved scanning tunneling microscopy, *Surf. Rev. Lett.* **25**, 1841003 (2018).
 - [180] M. Elbracht and M. Potthoff, Accessing long timescales in the relaxation dynamics of spins coupled to a conduction-electron system using absorbing boundary conditions, *Phys. Rev. B* **102**, 115434 (2020).
 - [181] R. M. Kaufmann, D. Li, and B. Wehefritz-Kaufmann, Notes on topological insulators, *Rev. Math. Phys.* **28**, 1630003.
 - [182] Y. Hancock, A. Uppstu, K. Saloriotta, A. Harju, and M. J. Puska, Generalized tight-binding transport model for graphene nanoribbon-based systems, *Phys. Rev. B* **81**, 245402 (2010).
 - [183] M. Z. Hasan and C. L. Kane, Topological insulators, *Rev. Mod. Phys.* **82**, 3045.
 - [184] L. Cano-Cortés, C. Ortix, and J. van den Brink, Fundamental differences between quantum spin Hall edge states at zigzag and armchair terminations of honeycomb and ruby nets, *Phys. Rev. Lett.* **111**, 146801 (2013).
 - [185] J. Kondo, Resistance minimum in dilute magnetic alloys, *Prog. Theor. Phys.* **32**, 37 (1964).
 - [186] C. Zener, Interaction between the d shells in the transition metals, *Phys. Rev.* **81**, 440 (1951).

-
- [187] H.-T. Elze, Linear dynamics of quantum-classical hybrids, *Phys. Rev. A* **85**, 052109 (2012).
 - [188] K.-H. Yang and J. O. Hirschfelder, Generalizations of classical Poisson brackets to include spin, *Phys. Rev. A* **22**, 1814 (1980).
 - [189] M. Sayad and M. Potthoff, Spin dynamics and relaxation in the classical-spin Kondo-impurity model beyond the Landau–Lifschitz–Gilbert equation, *New J. Phys.* **17**, 113058 (2015).
 - [190] C. Stahl and M. Potthoff, Anomalous spin precession under a geometrical torque, *Phys. Rev. Lett.* **119**, 227203 (2017).
 - [191] G. Lindblad, On the generators of quantum dynamical semigroups, *Commun. Math. Phys.* **48**, 119 (1976).
 - [192] P. Pearle, Simple derivation of the Lindblad equation, *Eur. J. Phys.* **33**, 805 (2012).
 - [193] J. Gao, W. Chen, X. C. Xie, and F.-c. Zhang, In-plane noncollinear exchange coupling mediated by helical edge states in quantum spin Hall systems, *Phys. Rev. B* **80**, 241302 (2009).
 - [194] A. Y. Kitaev, Fault-tolerant quantum computation by anyons, *Ann. Phys.* **303**, 2 (2003).
 - [195] G. Moore and N. Read, Non-Abelions in the fractional quantum Hall effect, *Nucl. Phys. B* **360**, 362 (1991).
 - [196] M. H. Freedman, A. Kitaev, M. J. Larsen, and Z. Wang, Topological quantum computation, *Bull. Amer. Math. Soc.* **40**, 31 (2003).
 - [197] C. Nayak and F. Wilczek, $2n$ -quasihole states realize $2n - 1$ -dimensional spinor braiding statistics in paired quantum Hall states, *Nucl. Phys. B* **479**, 529 (1996).
 - [198] M. H. Freedman, P/NP , and the quantum field computer, *Proc. Natl. Acad. Sci. U.S.A.* **95**, 98 (1998).
 - [199] S. Das Sarma, M. Freedman, and C. Nayak, Topologically protected qubits from a possible non-Abelian fractional quantum Hall state, *Phys. Rev. Lett.* **94**, 166802 (2005).
 - [200] P. Bonderson, M. Freedman, and C. Nayak, Measurement-only topological quantum computation, *Phys. Rev. Lett.* **101**, 010501 (2008).
 - [201] J. K. Pachos, *Introduction to Topological Quantum Computation*, (Cambridge University Press, 2012).
 - [202] S. Das Sarma, M. Freedman, and C. Nayak, Topological quantum computation, *Phys. Today* **59**, 32 (2006).
 - [203] G. P. Collins, Computing with quantum knots, *Sci. Am. Int. Ed.* **294**, 56 (2006), available at URL.
 - [204] E. Witten, Quantum field theory and the Jones polynomial, *Commun. Math. Phys.* **121**, 351–399 (1989).
 - [205] J. M. Leinaas and J. Myrheim, On the theory of identical particles, *Il Nuovo Cimento B* **37**, 1 (1977).
 - [206] A. Y. Kitaev, Unpaired Majorana fermions in quantum wires, *Phys. Usp.* **44**, 131 (2001).
 - [207] N. Read and D. Green, Paired states of fermions in two dimensions with breaking of parity and time-reversal symmetries and the fractional quantum Hall effect, *Phys. Rev. B* **61**, 10267 (2000).
 - [208] L. Fu and C. L. Kane, Superconducting proximity effect and Majorana fermions at the surface of a topological insulator, *Phys. Rev. Lett.* **100**, 096407 (2008).

-
- [209] J. Alicea, New directions in the pursuit of Majorana fermions in solid state systems, *Rep. Prog. Phys.* **75**, 076501 (2012).
 - [210] S. Tewari, S. Das Sarma, C. Nayak, C. Zhang, and P. Zoller, Quantum computation using vortices and Majorana zero modes of a $p_x + ip_y$ superfluid of fermionic cold atoms, *Phys. Rev. Lett.* **98** (2007).
 - [211] T. Sanno, S. Miyazaki, T. Mizushima, and S. Fujimoto, Ab initio simulation of non-Abelian braiding statistics in topological superconductors, *Phys. Rev. B* **103**, 054504 (2021).
 - [212] Y. Oreg, G. Refael, and F. von Oppen, Helical liquids and Majorana bound states in quantum wires, *Phys. Rev. Lett.* **105** (2010).
 - [213] F. Pientka, L. I. Glazman, and F. von Oppen, Topological superconducting phase in helical Shiba chains, *Physical Review B* **88** (2013).
 - [214] B. I. Halperin, Y. Oreg, A. Stern, G. Refael, J. Alicea, and F. von Oppen, Adiabatic manipulations of Majorana fermions in a three-dimensional network of quantum wires, *Phys. Rev. B* **85**, 144501 (2012).
 - [215] C. V. Kraus, P. Zoller, and M. A. Baranov, Braiding of atomic Majorana fermions in wire networks and implementation of the Deutsch–Jozsa algorithm, *Phys. Rev. Lett.* **111**, 203001 (2013).
 - [216] C. S. Amorim, K. Ebihara, A. Yamakage, Y. Tanaka, and M. Sato, Majorana braiding dynamics in nanowires, *Phys. Rev. B* **91**, 174305 (2015).
 - [217] M. Sekania, S. Plugge, M. Greiter, R. Thomale, and P. Schmitteckert, Braiding errors in interacting Majorana quantum wires, *Phys. Rev. B* **96**, 094307 (2017).
 - [218] D. Aasen, M. Hell, R. V. Mishmash, A. Higginbotham, J. Danon, M. Leijnse, T. S. Jespersen, J. A. Folk, C. M. Marcus, K. Flensberg, and J. Alicea, Milestones toward Majorana-based quantum computing, *Phys. Rev. X* **6**, 031016 (2016).
 - [219] T. Karzig, C. Knapp, R. M. Lutchyn, P. Bonderson, M. B. Hastings, C. Nayak, J. Alicea, K. Flensberg, S. Plugge, Y. Oreg, C. M. Marcus, and M. H. Freedman, Scalable designs for quasiparticle-poisoning-protected topological quantum computation with Majorana zero modes, *Phys. Rev. B* **95**, 235305 (2017).
 - [220] F. Harper, A. Pushp, and R. Roy, Majorana braiding in realistic nanowire Y-junctions and tuning forks, *Phys. Rev. Res.* **1**, 033207 (2019).
 - [221] C. Tutschku, R. W. Reinthaler, C. Lei, A. H. MacDonald, and E. M. Hankiewicz, Majorana-based quantum computing in nanowire devices, *Phys. Rev. B* **102**, 125407 (2020).
 - [222] Y. Tanaka, T. Sanno, T. Mizushima, and S. Fujimoto, Manipulation of Majorana-kramers qubit and its tolerance in time-reversal invariant topological superconductor, *Phys. Rev. B* **106**, 014522 (2022).
 - [223] K. Flensberg, F. von Oppen, and A. Stern., Engineered platforms for topological superconductivity and Majorana zero modes, *Nat. Rev. Mater.* **6**, 944 (2021).
 - [224] J. Li, H. Chen, I. K. Drozdov, A. Yazdani, B. A. Bernevig, and A. H. MacDonald, Topological superconductivity induced by ferromagnetic metal chains, *Phys. Rev. B* **90**, 235433 (2014).
 - [225] S. Nadj-Perge, I. K. Drozdov, J. Li, H. Chen, S. Jeon, J. Seo, A. H. MacDonald, B. A. Bernevig, and A. Yazdani, Observation of Majorana fermions in ferromagnetic atomic chains on a superconductor, *Science* **346**, 602–607 (2014).
 - [226] M. Ruby, F. Pientka, Y. Peng, F. von Oppen, B. W. Heinrich, and K. J. Franke, End states and subgap structure in proximity-coupled chains of magnetic adatoms, *Phys. Rev. Lett.* **115**, 197204 (2015).

-
- [227] S. Jeon, Y. Xie, J. Li, Z. Wang, B. A. Bernevig, and A. Yazdani, Distinguishing a Majorana zero mode using spin-resolved measurements, *Science* **358**, 772 (2017).
 - [228] R. Pawlak, M. Kisiel, J. Klinovaja, T. Meier, S. Kawai, T. Glatzel, D. Loss, and E. Meyer, Probing atomic structure and Majorana wavefunctions in mono-atomic Fe chains on superconducting Pb surface, *npj Quantum Inf.* **2** (2016).
 - [229] P. Krogstrup, N. L. B. Ziino, W. Chang, S. M. Albrecht, M. H. Madsen, E. Johnson, J. Nygård, C. Marcus, and T. S. Jespersen, Epitaxy of semiconductor-superconductor nanowires, *Nat. Mater.* **14**, 400–406 (2015).
 - [230] W. Chang, S. M. Albrecht, T. S. Jespersen, F. Kuemmeth, P. Krogstrup, J. Nygård, and C. M. Marcus, Hard gap in epitaxial semiconductor-superconductor nanowires, *Nat. Nanotechnol.* **10**, 232–236 (2015).
 - [231] J. E. Sestoft, T. Kanne, A. N. Gejl, M. von Soosten, J. S. Yodh, D. Sherman, B. Tarasinski, M. Wimmer, E. Johnson, M. Deng, J. Nygård, T. S. Jespersen, C. M. Marcus, and P. Krogstrup, Engineering hybrid epitaxial InAsSb/Al nanowires for stronger topological protection, *Phys. Rev. Mater.* **2**, 044202 (2018).
 - [232] J. Li, T. Neupert, B. A. Bernevig, and A. Yazdani, Manipulating Majorana zero modes on atomic rings with an external magnetic field, *Nat. Commun.* **7**, 10395 (2016).
 - [233] A. Wieckowski, M. Mierzejewski, and M. Kupczynski, Majorana phase gate based on the geometric phase, *Phys. Rev. B* **101**, 014504 (2020).
 - [234] T. Karzig, G. Refael, and F. von Oppen, Boosting Majorana zero modes, *Phys. Rev. X* **3**, 041017 (2013).
 - [235] W. Chen, J. Wang, Y. Wu, J. Qi, J. Liu, and X. C. Xie, Non-Abelian statistics of Majorana zero modes in the presence of an Andreev bound state, *Phys. Rev. B* **105**, 054507 (2022).
 - [236] T. Karzig, F. Pientka, G. Refael, and F. von Oppen, Shortcuts to non-Abelian braiding, *Phys. Rev. B* **91**, 201102 (2015).
 - [237] C. Knapp, M. Zaletel, D. E. Liu, M. Cheng, P. Bonderson, and C. Nayak, The nature and correction of diabatic errors in anyon braiding, *Phys. Rev. X* **6**, 041003 (2016).
 - [238] M. S. Scheurer and A. Shnirman, Nonadiabatic processes in Majorana qubit systems, *Phys. Rev. B* **88**, 064515 (2013).
 - [239] A. Conlon, D. Pellegrino, J. K. Slingerland, S. Dooley, and G. Kells, Error generation and propagation in Majorana-based topological qubits, *Phys. Rev. B* **100**, 134307 (2019).
 - [240] S. Bravyi and R. König, Disorder-assisted error correction in Majorana chains, *Commun. Math. Phys.* **316**, 641 (2012).
 - [241] B. Bauer, T. Karzig, R. V. Mishmash, A. E. Antipov, and J. Alicea, Dynamics of Majorana-based qubits operated with an array of tunable gates, *SciPost Phys.* **5**, 004 (2018).
 - [242] B. P. Truong, K. Agarwal, and T. Pereg-Barnea, Optimizing the transport of Majorana zero modes in one-dimensional topological superconductors, *Phys. Rev. B* **107**, 104516 (2023).
 - [243] P. Bonderson, M. Freedman, and C. Nayak, Measurement-only topological quantum computation via anyonic interferometry, *Ann. Phys.* **324**, 787 (2009).
 - [244] B. van Heck, A. R. Akhmerov, F. Hassler, M. Burrello, and C. W. J. Beenakker, Coulomb-assisted braiding of Majorana fermions in a Josephson junction array, *New J. Phys.* **14**, 035019 (2012).

-
- [245] M. Burrello, B. van Heck, and A. R. Akhmerov, Braiding of non-Abelian anyons using pairwise interactions, *Phys. Rev. A* **87**, 022343 (2013).
 - [246] T. Hyart, B. van Heck, I. C. Fulga, M. Burrello, A. R. Akhmerov, and C. W. J. Beenakker, Flux-controlled quantum computation with Majorana fermions, *Phys. Rev. B* **88**, 035121 (2013).
 - [247] L. Fu and C. L. Kane, Superconducting proximity effect and Majorana fermions at the surface of a topological insulator, *Phys. Rev. Lett.* **100**, 096407 (2008).
 - [248] J. C. Y. Teo and C. L. Kane, Topological defects and gapless modes in insulators and superconductors, *Phys. Rev. B* **82**, 115120 (2010).
 - [249] S. Nadj-Perge, I. K. Drozdov, B. A. Bernevig, and A. Yazdani, Proposal for realizing Majorana fermions in chains of magnetic atoms on a superconductor, *Phys. Rev. B* **88**, 020407 (2013).
 - [250] B. Braunecker, G. I. Japaridze, J. Klinovaja, and D. Loss, Spin-selective Peierls transition in interacting one-dimensional conductors with spin-orbit interaction, *Phys. Rev. B* **82**, 045127 (2010).
 - [251] T.-P. Choy, J. M. Edge, A. R. Akhmerov, and C. W. J. Beenakker, Majorana fermions emerging from magnetic nanoparticles on a superconductor without spin-orbit coupling, *Phys. Rev. B* **84**, 195442 (2011).
 - [252] M. Kjaergaard, K. Wölms, and K. Flensberg, Majorana fermions in superconducting nanowires without spin-orbit coupling, *Phys. Rev. B* **85**, 020503 (2012).
 - [253] I. Martin and A. F. Morpurgo, Majorana fermions in superconducting helical magnets, *Phys. Rev. B* **85**, 144505 (2012).
 - [254] F. Wilczek and A. Zee, Appearance of gauge structure in simple dynamical systems, *Phys. Rev. Lett.* **52**, 2111 (1984).
 - [255] E. C. Rowell and Z. Wang, Mathematics of topological quantum computing, *Bull. Amer. Math. Soc.* **55**, 183 (2018).
 - [256] J. C. Y. Teo and C. L. Kane, Majorana fermions and non-Abelian statistics in three dimensions, *Phys. Rev. Lett.* **104**, 046401 (2010).
 - [257] S. Rachel, Interacting topological insulators: a review, *Rep. Prog. Phys.* **81**, 116501 (2018).
 - [258] D. Krüger and M. Potthoff, Interacting Chern insulator in infinite spatial dimensions, *Phys. Rev. Lett.* **126**, 196401 (2021).
 - [259] A. Allerdt, A. E. Feiguin, and G. B. Martins, Spatial structure of correlations around a quantum impurity at the edge of a two-dimensional topological insulator, *Phys. Rev. B* **96**, 035109 (2017).
 - [260] A. Allerdt, A. Feiguin, and G. Martins, Kondo effect in a two-dimensional topological insulator: Exact results for adatom impurities, *J. Phys. Chem. Solids* **128**, 202 (2019).
 - [261] N. Lenzing, A. I. Lichtenstein, and M. Potthoff, Emergent non-Abelian gauge theory in coupled spin-electron dynamics, *Phys. Rev. B* **106**, 094433 (2022).
 - [262] J. C. Baez, The tenfold way, arXiv:2011.14234.
 - [263] J. Baez, The tenfold way: 101 topics in contemporary mathematics (2025), available at URL.
 - [264] G. Refael, Majorana fermions (2015), lecture notes from the 15th ASC School on Topological Matter, LMU Munich, available at URL.

C – List of Publications

- [RQ1] S. Michel, A. Fünfhaus, R. Quade, R. Valentí, and M. Potthoff, Bound states and local topological phase diagram of classical impurity spins coupled to a Chern insulator, *Phys. Rev. B* **109**, 155116 (2024).
- [RQ2] R. Quade and M. Potthoff, Controlling the real-time dynamics of a spin coupled to the helical edge states of the Kane–Mele model, *Phys. Rev. B* **105**, 035406 (2022).
- [RQ3] R. Quade and M. Potthoff, Exchangeless braiding of Majorana zero modes in weakly coupled Kitaev chains, *Phys. Rev. B* **112**, 085411 (2025).

Declaration of Contributions

Publication [RQ1]: The author is co-author of this publication. They contributed to the development of the topological aspects of the problem and helped to interpret the results in the context of the topological phases of matter. In addition, they engaged in scientific discussions and assisted in the publication process.

Publication [RQ2]: The author is the main author of this publication. They developed the code, performed the numerical computations, generated the figures, and interpreted the results. The project was completed under the supervision of Prof. Dr. Potthoff, whose support is acknowledged with gratitude.

Publication [RQ3]: The author is the main author of this publication. They developed the code, performed the analytical and numerical computations, generated the figures, and interpreted the results. The project was completed under the supervision of Prof. Dr. Potthoff, whose support is acknowledged with gratitude.

This thesis was typeset using L^AT_EX. Figures were created with the Matplotlib library for Python and the vector graphics editor Inkscape. Language was refined with the help of ChatGPT (OpenAI) and DeepL. The pink doughnut on the title page was generated using ChatGPT.

Eidesstattliche Versicherung

Hiermit versichere ich an Eides statt, die vorliegende Dissertationsschrift selbst verfasst und keine anderen als die angegebenen Hilfsmittel und Quellen benutzt zu haben.

Hamburg, den _____

Robin Leon Quade