



UNIVERSIDADE ESTADUAL PAULISTA
“JÚLIO DE MESQUITA FILHO”
Câmpus de São José do Rio Preto

Pedro Henrique Muller Bortolucci

Introduction to Quantum Invariants of Knots and a Path
Towards Topological Quantum Field Theories (TQFT)

São José do Rio Preto
2023

Pedro Henrique Muller Bortolucci

Introduction to Quantum Invariants of Knots and a Path
Towards Topological Quantum Field Theories (TQFT)

Dissertação apresentada como parte dos requisitos para obtenção do título de Mestre em Matemática, junto ao Programa de Pós-Graduação em Matemática, do Instituto de Biociências, Letras e Ciências Exatas da Universidade Estadual Paulista “Júlio de Mesquita Filho”, Câmpus de São José do Rio Preto.

Orientadora: Profa. Dra. Alice Kimie Miwa Libardi

Co-orientador: Prof. Dr. Louis Funar

Financiadora: CAPES

São José do Rio Preto
2023

B739i

Bortolucci, Pedro Henrique Muller

Introduction to Quantum Invariants of Knots and a Path Towards
Topological Quantum Field Theories (TQFT) / Pedro Henrique Muller

Bortolucci. -- São José do Rio Preto, 2023

208 p. : il.

Dissertação (mestrado) - Universidade Estadual Paulista (Unesp), Instituto
de Biociências Letras e Ciências Exatas, São José do Rio Preto

Orientadora: Alice Kimie Miwa Libardi

Coorientador: Louis Funar

1. Topologia Algébrica. 2. Variedades. 3. Nós. 4. Invariantes Quânticos. I.
Título.

Pedro Henrique Muller Bortolucci

Introduction to Quantum Invariants of Knots and a Path
Towards Topological Quantum Field Theories (TQFT)

Dissertação apresentada como parte dos requisitos para obtenção do título de Mestre em Matemática, junto ao Programa de Pós-Graduação em Matemática, do Instituto de Biociências, Letras e Ciências Exatas da Universidade Estadual Paulista “Júlio de Mesquita Filho”, Câmpus de São José do Rio Preto.

Financiadora: CAPES

Comissão Examinadora

Prof. Dr. Louis Funar- Coorientador
Institut Fourier - Université Grenoble Alpes

Prof. Dr. Thiago de Melo
UNESP - Rio Claro

Prof. Dr. Luiz Roberto Hartmann
UFSCar - São Carlos

Rio Claro
07 de fevereiro de 2023

AGRADECIMENTOS

Gostaria de agradecer minha orientadora Alice Kimie Miwa Libardi pela imensa ajuda, suporte e incentivo que me deu ao longo de 5 anos trabalhando ao seu lado. À professora Renata Zotin Gomes por incontáveis vezes que me ajudou ao longo de minha graduação e pós graduação. Ao professor Louis Funar que me recebeu em seu país e me forneceu uma das experiências mais ricas de toda a minha vida.

Agradeço também à todas as instituições que me forneceram incrível suporte ao longo desses anos, UNESP campi de Rio Claro e de São José do Rio Preto e também ao Institut Fourier, onde passei o período de 2 meses estudando.

Pelo auxílio financeiro recebido ao longo de minha pós graduação, agradeço à CAPES.

Finalmente, não posso deixar de agradecer à minha família, namorada e amigos, pelo suporte, incentivo e palavras de carinho que foram fundamentais para mim tanto em momentos de dificuldade quanto em momentos de prazer.

O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior - Brasil (CAPES) - Código de Financiamento 001

RESUMO

Os invariantes quânticos surgem por volta do final dos anos 1980 com a descoberta do polinômio de Jones e se desenvolveu como uma nova linha de pesquisa em matemática, mais geralmente chamada de Topologia Quântica. De fato, esses invariantes quânticos são invariantes topológicos de grande relevância no estudo da topologia algébrica. O objetivo desse trabalho é estudar os invariantes quânticos que permitam uma completa classificação dos nós, além da introdução de alguns conceitos como os de álgebras de Lie, álgebras de Hopf e de teorias topológicas de campos quânticos (TQFT), com o intuito de preparar um estudo mais aprofundado nessa área.

O trabalho será baseado no pre-print: "A Brief Introduction to Knot Invariants" de Louis Funar, bem como em trabalhos clássicos de TQFT's por Turaev, como o livro "Quantum Invariants of Knots and 3-Manifolds".

Palavras-chave: Topologia Algébrica. Variedades. Nós. Invariantes Quânticos.

ABSTRACT

Quantum invariants emerged in the late 1980s with the discovery of the Jones polynomial and developed as a new line of research in mathematics, more often called Quantum Topology. In fact, it is a topological invariant, of relevance in the classification in Algebraic Topology. The objective of this work is to study the invariants which allow a complete classification of the knots, as well as the introduction of some concepts such as Lie Algebras, Hopf Algebras and the Topological Quantum Field Theories, with the objective of preparing a more deep study of problems in the area.

The work will be based on the pre-print: "A Brief Introduction to Knot Invariants" by Louis Funar, as well as classical works on TQFT's by Turaev, such as the book "Quantum Invariants of Knots and 3-Manifolds".

Keywords: Algebraic Topology. Manifolds. Knots. Quantum Invariants.

List of Figures

2.1	(<i>a</i>) is a good representation and (<i>b</i>) is not a good representation	24
2.2	The trefoil and the figure eight.	25
2.3	Triangular move	26
2.4	5 different ways of triangular moves, with minimal number of intersections.	26
2.5	Different kinds of Reidemeister moves.	27
2.6	An unknot diagram	28
2.7	Connected sum of a trefoil (3_1) and a figure eight knot (4_1).	31
2.8	Barycentric subdivision	32
2.9	Rotation	32
2.10	Types of crossings	33
2.11	Linking number of R_2 move	34
2.12	Linking number of R_3 move	34
2.13	A first try on constructing a Seifert surface.	37
2.14	The Seifert surface of the Figure Eight knot.	38
2.15	The Seifert surface of the Trefoil knot.	39
3.1	Table of Jones Polynomials	52
3.2	Two mutant knots	53
3.3	A non-trivial link, with trivial Jones polynomial	54
3.4	Prime, non prime, split, and non split links	55
3.5	Prime and strongly prime links	56
3.6	A reducible crossing	56
3.7	Another way of identify a reducible crossing	57
3.8	Smoothings of a state	58
3.9	Smoothing of a state, with the notation of labels	58
3.10	State on the trefoil T	58
4.1	Operation on the braid group	68
4.2	The closure of a braid	69
4.3	The Folding Transformation of a Tree.	72
4.4	A Movement of Type- $\mathcal{E}$	74
4.5	A Movement of Type- $\mathcal{S}$	76
4.6	Division of space into 4 regions.	77
5.1	A braid that represents the figure eight knot.	101
5.2	A method for calculating with diagrams.	135
6.1	Coxeter graphs of $A_1 \times A_1$, A_2 , B_2 and G_2	149
6.2	Dynkin diagrams of B_2 and G_2	149

7.1	A cobordism between S^1 and $S^1 \amalg S^1$	170
7.2	The gluing of two cobordisms.	171
7.3	How to obtain $(M, \emptyset, S^1 \amalg S^1 \amalg S^1)$ from $(M, S^1, S^1 \amalg S^1)$	172

Contents

Introduction	15
1 Preliminaries	17
1.1 Some Differential Topology Language	17
1.2 The PL Topology	18
1.3 Problems with triangulations: The Hauptvermutung	21
1.4 Embeddings and Knottedness	22
2 Knots and Links	23
2.1 Initial Concepts	23
2.2 Diagrams of Knots and Links	24
2.3 Knot Complements	28
2.4 Triangulations on Surfaces	31
2.5 The Linking Number	33
2.6 Seifert Surfaces	36
2.7 Milnor's Invariants or The Higher Linking Numbers	41
3 The First Quantum Invariants	45
3.1 The Kauffman Bracket	45
3.2 The Writhe and the Jones Polynomial	47
3.3 Properties of the Jones Polynomial	50
3.4 Names of knots and the Jones polynomial table	51
3.5 Mutation	53
3.6 Alternating Links	54
3.7 The Tait Conjectures	57
3.8 The writhe of alternating links	64
4 The Braid Groups	67
4.1 Some group notation	67
4.2 A geometric interpretation	68
4.3 Equivalence of definitions	68
4.4 Alexander's theorem	69
4.5 The Vogel's Algorithm	71
4.6 Markov's Theorem	74
4.6.1 Initial Concepts	74
4.6.2 Useful Lemmas for Proving the Markov's Theorem	78
4.7 Another Version of Markov's Theorem	91

5	Quantum Link Invariants	97
5.1	Markov's Traces	97
5.2	The Burau Representation	98
5.3	Reduced Burau Representation	99
5.4	Some Consequences of Studying Representations of B_n	101
5.4.1	Hecke Algebras	102
5.4.2	Birman-Murakami-Wenzl Algebras	103
5.5	The HOMFLY Polynomial	103
5.6	HOMFLY polynomial - Properties	105
5.7	The Kauffman Polynomial	111
5.8	The Yang-Baxter Equation	121
5.9	The G_2 Invariant	125
6	Lie Algebras	137
6.1	The Definition and The Classical Lie Algebras	137
6.1.1	The Classical Algebras: A_ℓ , B_ℓ , C_ℓ and D_ℓ	138
6.1.2	Representations of Lie Algebras	139
6.2	Classification of Semisimple Lie Algebras	139
6.2.1	Nilpotent Lie Algebras	139
6.2.2	Casimir Element of a Representation	140
6.2.3	Toral Subalgebras	141
6.2.4	Root Systems	142
6.2.5	Cartan Matrix	147
6.2.6	Dyinkin Diagrams	149
7	Topological Quantum Field Theories	151
7.1	Monoidal Categories	151
7.1.1	Braiding and Twists in Monoidal Categories	152
7.1.2	Duality in Monoidal Categories	154
7.2	Ribbon Categories	154
7.2.1	Traces and Dimensions	154
7.2.2	A Visual Study of Ribbon Categories	155
7.3	Ribbon Graph Category	159
7.4	Modular Tensor Categories	165
7.5	Invariants of 3-Manifolds	165
7.6	The definition of a TQFT	167
7.6.1	Cobordisms of $\mathfrak{A}$ -spaces	170
7.6.2	Definition of Topological Quantum Field Theory	171
	Referências	175
A	A Review on Homology and Cohomology	177
A.1	General Review	177
A.2	Low Dimensional Manifolds	199

Introduction

The history of topology, begins with the brilliant mathematician Leonhard Euler in 1736 with the solution to a famous problem called *The Seven Bridges of Königsburg*, Euler initiated the studies on graph theory, which later became a big branch of topology. Euler noticed that there were some problems that were not exactly geometry problems, but they indeed looked like it, hence Euler called those kinds of studies "geometry of position". Later on Henri Poincaré a great topologist with enormous contributions to algebraic and differential topology, separated the study of "geometry of position" into a new branch of mathematics, which he called "Analysis Situs" and later on, it became what is today called *Topology*.

The interest in knots, began with physicists, before the discoveries of the structure of the fundamental part of matter, the atoms, by Ernest Rutherford. More precisely, Lord Kelvin thought that the atoms were actually really small loops knotted in such a way that different kinds of knottedness would result in different chemical elements. But, with the discovery of the real atomic structure, the physicists lost interest in knots, while mathematicians found that incredibly interesting and continued to study the properties and classification of knots. The main name that appears on the beginning of mathematical studies of knots is Carl Friedrich Gauss, with the discovery of the Gauss linking integral in 1833, giving a way to compute the linking number of a link, by calculating an integral.

To help the process of classification of knots, various concepts and topological invariants were created, concepts such as the Reidemeister moves, the linking number, the Alexander Polynomial and the knot group helped mathematicians to distinguish different knots. Later on, Vaughan Jones discovered one of the most powerful tools in the study of knots, the Jones Polynomial, although there are still open problems involving the Jones polynomial, such as the unproven statement saying that if a knot has Jones polynomial equal to 1, then it is the unknot, the Jones polynomial helped develop the beautiful theory of knots.

Today, the knot theory, viewed from a mathematical perspective is a still young theory, with some really interesting open problems and with lots of discoveries still to be made. From an application point of view, the knot theory has helped scientists study untangling of DNA molecules, as well as it has played a role on the study of various topics in physics, for example topological fluid dynamics. And maybe, who knows, Gauss were actually kind of right, if we think of the basic elements of string theory, small loops that vibrate generating different kinds of forces, such as the gravitational force. Maybe, the knot theory is even related as the theory of everything. In the first chapter we introduce some notion of PL topology, which will be implicitly used in the rest of the text, since the links diagrams that we will be studied can be seen to have a PL structure. Chapters 2 and 3 are dedicated to the study of the main definitions of knot theory as well as some

useful results. Chapters 4 and 5 are dedicated to making connections between the link invariants defined on the previous chapters with the study of braids. Also we prove the Markov's theorem, which is important in defining new linking invariants, finally we end these chapters with a generalization of link projections, which now can be assumed to have triple points.

The rest of the dissertation will focus on a review of the main results in homology theory and its applications in manifolds of low dimensional, i.e., dimensions 2, 3 and 4. We also define concepts important in the study of such objects, namely the intersection forms, the trilinear forms on 3-manifolds and the linking forms defined on the torsion part of the homology group of manifolds. We also introduce the main definitions and some results in Lie Algebras, more precisely those necessary to get to the definition of the Dynkin diagrams and the classification theorem for simple Lie algebras. And we finish our work defining, in a very categorical point of view, topological quantum field theory.

It is important to note that chapters 6, 7 and 8 are more focused in introducing concepts rather than proving them. Although we will prove some results, the main objective is to discuss and understand what are some key concepts in the Lie algebra theory and the TQFT's theory that may possibly be used in a future doctorate's thesis.

1 Preliminaries

1.1 Some Differential Topology Language

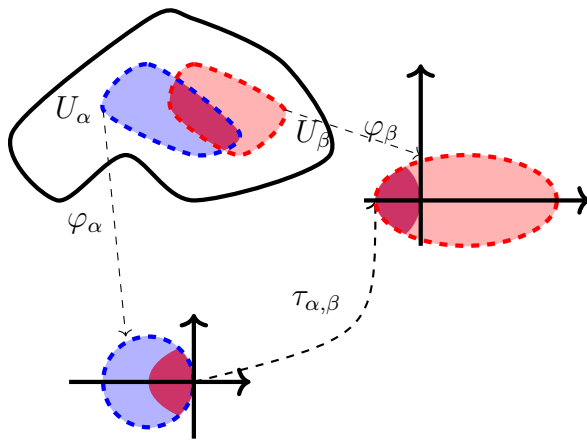
First we are going to introduce some concepts on manifolds topology of any dimension so we can start our studies on lower dimensional manifolds, meaning manifolds of 4 or lower dimensions. Actually, our studies will have great focus on *links*, that can, in general, be thought of as disjoint unions of loops in $\mathbb{R}^3$, that can have only one component and in this case the link will be called a *knot*.

In fact it is common to think on links as embedding of S^1 in S^3 and as a mean of visualization, we can think of S^3 as the *one point compactification* of $\mathbb{R}^3$. It is a simple general topology exercise to prove that $\mathbb{R}^3 \cup \{\infty\}$ is a compact space as well as show that $S^3 \simeq \mathbb{R}^3 \cup \{\infty\}$. Hence, working with links in S^3 can be simplified as working with links in $\mathbb{R}^3$.

Now, we will introduce some concepts and results (not all of such will be proved) that will be useful later on.

Definition 1.1.1. A Hausdorff separable topological space M is a *Topological Manifold* if each point in M has a open neighborhood homeomorphic to an open subset of $\mathbb{H}^n$, where $\mathbb{H}^n = \{(x_1, \dots, x_n) \in \mathbb{R}^n; x_n \geq 0\}$. Such open neighborhoods are called *Charts* and a open cover of charts is called an *Atlas*.

It is common to denote the charts on a manifold as $(U_\alpha, \varphi_\alpha)$, for example, if we want to say that U_α is an open set in a topological manifold M , that is homeomorphic to an open subset of $\mathbb{H}^n$ via φ_α . Thus, we have the following definitions:



Definition 1.1.2. Let $(U_\alpha, \varphi_\alpha)$ and (U_β, φ_β) be two charts on a topological manifold M , such that $U_\alpha \cap U_\beta \neq \emptyset$, we have that the map $\tau_{\alpha,\beta}: \varphi_\alpha(U_\alpha \cap U_\beta) \rightarrow \varphi_\beta(U_\alpha \cap U_\beta)$, defined as $\tau_{\alpha,\beta} = \varphi_\beta \circ \varphi_\alpha^{-1}$ is called *Transition Map*.

Definition 1.1.3. A topological manifold is a C^r -manifold if it contains an atlas such that the transition maps between the images of charts that intersect in $\mathbb{H}^n$ are maps of class C^r .

Figure 1.1: The transition map

Let us now remember that two maps, f and g , are called *Isotopic by Homeomorphisms* if they are homotopic via a homotopy H and such that each map $H(t, x)$ is a homeomorphism. Furthermore, recall that a map between open sets of $\mathbb{R}^m$ is a *Diffeomorphism* if, and only if, it is a differentiable homeomorphism with differentiable inverse.

Then, we will say two C^r -structures on a manifold are *Isotopic*, when the identity map is isotopic via homeomorphism to a C^r -diffeomorphism.

The concepts of C^r -manifolds exists for $r \in \{1, \dots, \infty\}$, that are the values where the concept of C^r class itself exists. However, the most interesting values for our studies will be the values in which the transition maps for a certain atlas of a manifold are continuous, but not necessarily have continuous derivatives, and the values for which the transition maps are infinitely differentiable and have continuous derivative of any order, in other words, will be interesting for us, the cases where $r = 0$ and $r = \infty$. The following Munkres result gives us the following:

Theorem 1.1.1. *For any $1 \leq r \leq k \leq \infty$ a C^r -manifold have a C^k -structure, for which the C^r -structure is C^r -homeomorphic to the initial C^k -structure. This C^k -structure is unique up to isotopy via C^r -diffeomorphisms.*

A simple way to see that is to simply note that for any C^k -atlas in the C^r -manifold M , we can consider the transition maps of class C^k to be of class C^r , since $r \leq k$, so the C^r -structure on M is induced by the C^k -structure and it is, by construction, compatible with the initial structure.

There are three principal types of relevant structures for the manifold topology, these types of structure are the TOP, PL and DIFF. The earlier discussion, allows us to define two of them, TOP and DIFF, the PL category will be discussed in detail later on.

Definition 1.1.4. TOP is the category of rough manifolds, in other words, the manifolds with a C^0 -structure. DIFF is the category of smooth manifolds, i.e., manifolds with a C^∞ -structure.

1.2 The PL Topology

There exists an intermediary category between those two previously mentioned, called the PL category (PL stands for Piecewise Linear) in which the transition maps for a given atlas on a topological manifold are piecewise linear maps with respect to the linear structure of $\mathbb{H}^n$. We will study in a certain depth this theory.

The PL topology, also called combinatorial topology, was the first kind of algebraic topology, it was born with the study of topologies of low dimension, facilitated by the fact that we can visualize a great part of its concepts. This kind of topological study uses simplicial complexes and triangulations, as a matter of fact, some big problems in this area are related to triangulation of manifolds like the *Hauptvermutung* problem, German for "the main conjecture", mainly studied in "*Einführung die Hauptvermutung der kombinatorischen Topologie*".

Because of some particular problems, like the fact that the dual of a simplex not necessarily be a simplex, the combinatorial topology, or PL topology, ended up by being, in a way, replaced by the algebraic topology, which possesses some extremely powerful tools for the study of topological invariants, however, this does not implies that combinatorial topology has become less interesting, actually, some applications of this kind of study appears on computer sciences for example.

We start our study looking to the concept of a *join* of two sets. Consider A and B two arbitrary sets, we define the join of A and B , denoted by AB , as the set

$$AB = \{\lambda a + \mu b; a \in A, b \in B, \lambda, \mu \in \mathbb{R}_+ \text{ e } \lambda + \mu = 1\}.$$

When one of these sets, say A , is a singleton $\{a\}$, any point on $\{a\}B$ is represented in a unique way, then we will call this join a *cone* and denote it simply by aB . The Figure 1.2, below, represents precisely one case in which we have a cone, and one case in which the join is not a cone, moreover, we can also have a visual idea of a join and a cone.

The definition of a cone, can be simplified by saying that for every two distinct points b_1 and b_2 on the *basis* B , the segments ab_1 and ab_2 intersects only on the point a . Then, it is easy to see in Figure 1.2 that the second picture is not a cone, since there are two distinct segments that intersect in more than one point.

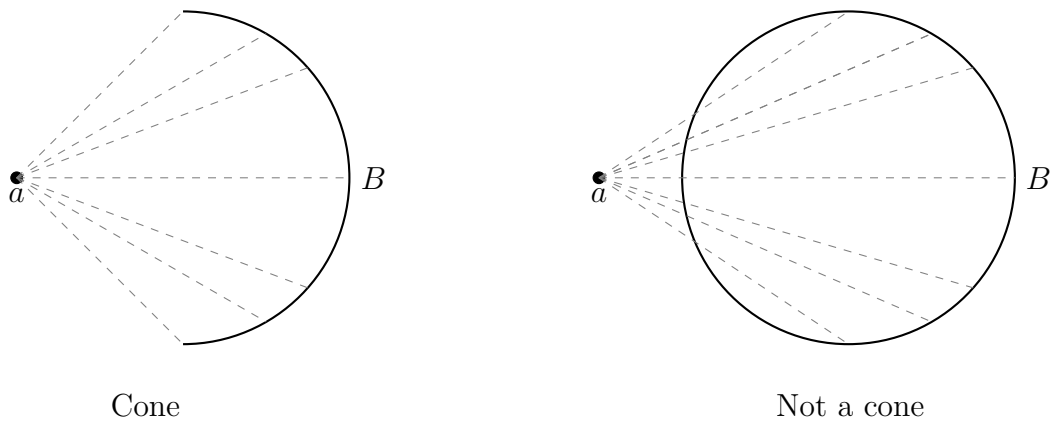


Figure 1.2: A cone and a join

Definition 1.2.1. A subset P of $\mathbb{R}^n$ is called a *polyhedron* if each point $a \in P$ has a conical neighborhood $N = aL$, where L is a compact set; N is also called *star* of a in P , and L is called a *link*, we denote then $N = N_a(P)$ and $L = L_a(P)$.

If we consider p any point in $\mathbb{R}^n$ and L a compact set such that pL is a cone, we can also restrict pL to a ε -neighborhood of p in the following way: since L is compact, then there exists $\varepsilon > 0$ sufficiently small, in such a way that $d(p, L) \geq \varepsilon$ (remember $\mathbb{R}^n$ is also a metric space), thus, let $N_\varepsilon(p, P) = \{x; x \in P, d(p, x) \leq \varepsilon\}$ and $\dot{N}_\varepsilon(p, P) = \{x; x \in P, d(p, x) = \varepsilon\}$ so $N_\varepsilon(p, P) = p\dot{N}_\varepsilon(p, P)$ is a cone as well as a ε -neighborhood of p .

As we said earlier, the use of simplexes is of great use in combinatorial topology, thus we introduce some of the concepts needed for this preliminaries chapter, after that, we will establish a theorem, which the proof will not be entirely given, but a good idea is presented. This theorem allows us to relate the study of polyhedrons with the study of simplexes, since (as we'll see) we can understand polyhedrons as locally the union of simplexes.

Definition 1.2.2. Consider a set $\{v_0, v_1, \dots, v_n\} \subset \mathbb{R}^n$, we say this set is *independent* if the set of vectors $\{v_1 - v_0, v_2 - v_0, \dots, v_n - v_0\}$ is linearly independent. The set $A \subset \mathbb{R}^n$ defined as the repeated join $v_0v_1v_2 \dots v_n$ is called a *n-simplex* of vertices v_i , and we say that this set of vertices *span* the *n-simplex* A . Finally, we say that the simplex, B , spanned by a subset of the vertices of A is a *face* of A and denoted by $B < A$.

Theorem 1.2.1. *A compact polyhedron is the finite union of simplexes. In general, a polyhedron is locally a finite union of simplexes.*

Proof. Let P be a polyhedron in $\mathbb{R}^m$ and define the subspace $\langle P \rangle \subset \mathbb{R}^m$ as the space generated by P as being the intersection of all subspaces V of $\mathbb{R}^m$ containing P .

We proof by induction on the dimension of $\langle P \rangle$, which we shall denote by n . Without loss of generality, suppose $P \subset \mathbb{R}^n$. Taking $a \in P$, we have that since P is a polyhedron, there is a star aL of a for some subset L of P . But as we noted before, we can work with an ε -neighborhood of a .

Now, let F be a *proper* face of $N_\varepsilon(a, \mathbb{R}^n)$, i.e., a face that is not equal the set $N_\varepsilon(a, \mathbb{R}^n)$. Then $\dim(F \cap L) < n$, thus $F \cap L$ is a finite union of simplexes. Since $L \subset N_\varepsilon(a, \mathbb{R}^m)$, it follows that L is a finite union of simplexes, say $L = \bigcup_i A_i$. Thus $aL = a \left(\bigcup_i A_i \right) = \bigcup_i aA_i$ is also a finite union. Hence for each point p in P , we have this property, so therefore $P = \bigcup_{p \in P} \left(\bigcup_i pA_i \right)$, but by hypothesis P is compact, there exists a finite set of open sets $B_j = \bigcup_i p_j A_i$, $j \in \{1, \dots, n\}$, such that $P = \bigcup_{j=1}^n (\bigcup_i p_j A_i)$, hence P is the finite union of simplexes. $\square$

We are now going to introduce the definition of a PL map, and how we can use this to see that the graph of a PL map between polyhedra is also a polyhedron. Then we introduce the concept of a subdivision and triangulation, two very important concepts in the PL topology.

Definition 1.2.3. A map $f: P \rightarrow Q$ between polyhedra is a *piecewise-linear* map, i.e., a PL map if for each point p we take in P there exists a star $N = pL$ such that $f(\lambda a + \mu x) = \lambda f(a) + \mu f(x)$ where x is in the compact L and $\lambda, \mu \in \mathbb{R}$ with $\lambda + \mu = 1$.

We can understand the concept of a PL map as a map that preserves the PL structure of a polyhedron P with respect to the PL structure of Q , as said in [4] " f is locally conical", meaning that each line connecting the base of a cone to its vertex in P is linearly sent to its image in Q .

Observe now that if f is a PL map between polyhedra, then denoting $\Gamma(f)$ the graph of f , i.e.,

$$\Gamma(f) = \{(x, f(x)) \in \mathbb{R}^{n+1}; x \in P\}$$

then $\Gamma(f)$ is also a polyhedron. In fact, if f is a PL map, then $f(\lambda x + \mu y) = \lambda f(x) + \mu f(y) \Leftrightarrow \lambda(x, f(x)) + \mu(y, f(y)) = (\lambda x + \mu y, \lambda f(x) + \mu f(y)) = (\lambda x + \mu y, f(\lambda x + \mu y)) = (z, f(z))$, for $z = \lambda x + \mu y \in P$. In other words, the points of $\Gamma(f)$ respect the conical structure, i.e., it is a polyhedron.

We can extend the concept of PL map to a *PL Homeomorphism*, by simply saying a map is a PL homeomorphism whenever it is a homeomorphism and it is PL as a map between polyhedra.

Definition 1.2.4. A *simplicial complex* $\mathcal{K}$ is a collection of simplexes satisfying the following conditions:

- i. If $C \in \mathcal{K}$ and B is a face of C , then $B \in \mathcal{K}$;
- ii. If B and C are two simplexes in $\mathcal{K}$, then $B \cap C$ is a face of both B and C .

As we have seen, any compact polyhedron can be seen as a finite union of simplexes. But if we take a simplicial complex we can also construct a polyhedron from it by just taking the union of all of the simplexes in the complex. This polyhedron is called *the underlying polyhedron* of the simplicial complex $\mathcal{K}$, and it is denoted by $|\mathcal{K}|$. Therefore, any polyhedron P define a simplicial complex $\mathcal{K}$, by just taking $\mathcal{K}$ as the simplicial complex in which we have $|\mathcal{K}| = P$.

Definition 1.2.5. L is a *subdivision* of a simplicial complex $\mathcal{K}$, when $|L| = |\mathcal{K}|$ and each simplex of L is contained in a simplex of $\mathcal{K}$. We denote it by $L \triangleleft \mathcal{K}$.

Definition 1.2.6. Let $a \in |\mathcal{K}|$, we say that $K' \triangleleft \mathcal{K}$ is obtained by *starring* in a if K' is obtained from $\mathcal{K}$ by substituting each simplex $C \in \mathcal{K}$ with $a \in C$ by the simplex aB , with $B = \{F; F < C, a \notin F\}$. The result of starring by the points $a_1, a_2, \dots, a_r \in |\mathcal{K}|$ in order, is called *stellar subdivision*.

Definition 1.2.7. Let K and L be simplicial complexes, then $f: |K| \rightarrow |L|$ is called a *simplicial map* if for each $C \in K$ $f|_C$ is linear and $f(C)$ is a simplex in L .

Now, we can say that a PL map induces a simplicial map between the simplicial complexes associated to the polyhedra. Moreover, we can say that two polyhedra are PL homeomorphic if and only if, each of their associated simplicial complexes are isomorphic after an adequate choice of subdivisions. Finally, we can talk about triangulations.

Definition 1.2.8. A *triangulation* of a topological space $\mathbb{X}$ is a simplicial complex for which each underlying polyhedron is homeomorphic to $\mathbb{X}$. A topological space $\mathbb{X}$ is said to be *triangulable* if it admits a triangulation.

1.3 Problems with triangulations: The Hauptvermutung

As well as the main objective of topology, in general, is to classify topological spaces up to homeomorphisms, the main objective of manifold topology is to classify topological manifolds as well as find a way to decide if a given manifold has a PL or a DIFF structure. The first result in this direction, proved in [18], is a consequence of a theorem called *The Simplicial Approximation Theorem*, and helps us to decide the structure of a topological manifold:

Theorem 1.3.1. *Every continuous map between polyhedra $|P|$ and $|Q|$ is homotopic to the topological realization of a simplicial map $P' \rightarrow Q$, where P' is a simplicial subdivision of P .*

Remember that $|P|$ is the underlying polyhedron of the simplicial complex P , and since P' is a subdivision of P , by definition, $|P'| = |P|$. In particular, what the theorem allows us to say is that every continuous map between polyhedra is homotopic to a PL map, namely the topological realization (or *simplicial approximation*, in the words of [18]) from theorem 1.3.1. But it leads us to a problem, which has been given a distinguished German name, the Hauptvermutung. Which roughly speaking is the generalization of what we have just stated. The first hauptvermutung is given as:

Polyhedral Hauptvermutung 1: Every homeomorphism of polyhedra is homotopic to a PL homeomorphism.

For low dimensions such as 1 and 2, Papakyriakopoulos has proven this Hauptvermutung. For 6 dimensional manifolds, Milnor gave the first counterexample for the first Hauptvermutung, meaning we have it valid for $n \leq 2$, but not valid for $n = 6$.

In dimension 3, we have equivalence between PL, TOP and DIFF structures, as given in the following theorem:

Theorem 1.3.2. *Any compact 3-manifold is triangulable and has an unique DIFF structure.*

The proof to this theorem can be found in [16] (page 50). The previous theorem allows to interchange terms such as diffeomorphism and homeomorphism when dealing with 3-dimensional manifolds.

1.4 Embeddings and Knottedness

Before actually start talking about knots, we must give some notation that are going to be useful throughout the text, we need to know what does it mean to say thing such as " M unknots in Q ", or what is an isotopy in the context of PL topology. Therefore, the following definitions fulfills our purpose.

Definition 1.4.1. Let Y be a finite simplicial complex. A *isotopy* of Y is a PL map $h: Y \times [0, 1] \rightarrow Y \times [0, 1]$ such that $h(Y \times \{t\}) = Y \times \{t\}$ and $h|_{Y \times \{t\}}$ is a PL-homeomorphism for each $t \in [0, 1]$ and $h_0 = Id_Y$.

Definition 1.4.2. If X and Y are two simplicial complexes and $f, g: X \rightarrow Y$ are embeddings, then f and g are isotopic if there exists an isotopy h on Y such that $h_1(\bullet, 1) \circ f = g$.

Remember that an embedding is a homeomorphism onto its image, meaning that a function $f: X \rightarrow Y$ is an embedding when f is injective and X is homeomorphic to $f(X)$. Now we can see what "unknots" mean in this context.

Definition 1.4.3. M unknots in Q if any two equivalent embeddings of M in Q are isotopic.

So, we can understand this as taking the space M , and embedding it onto Q in many different ways, but every time we take two different embeddings, we can find a smooth deformation by homeomorphisms that connects one embedding to the other, i.e., an isotopy between the two embeddings.

Definition 1.4.4. A link L is a *tame link* if it consists of disjoint simple curves made from a finite number of linear segments.

We usually work with links in S^3 , but only with the purpose of visualization, we can think of S^3 as the space $\mathbb{R}^3$ with a point at infinity, or as the one point compactification of $\mathbb{R}^3$. Since we work in S^3 , the linear structure is that of the boundary of a 4-simplex.

2 Knots and Links

2.1 Initial Concepts

We are now able to begin the study of knots and links. First let's define what it mean to say that two knots are equivalent in a mathematical way. Intuitively we can see that two knots in $\mathbb{R}^3$ are equivalent, by the following method: suppose you have two apparently different knots K_1 and K_2 , when you are able to deform K_1 , without cutting it or passing it through itself, into K_2 , then visually we can see that K_1 and K_2 are "the same" or equivalent. Or in the same sense, every deformation of that kind (no cuts or passing through itself) on a given knot, gives a new knot that is equivalent to the given one.

Now, in terms of isotopies we can define this equivalence as follows:

Definition 2.1.1. Two links are *equivalent* if there exists a PL homeomorphisms that preserves orientation of S^3 , transforming one link into the other. If the two links are oriented, then the orientation is preserved under this equivalence.

The last sentence on the definition means that if the orientation preserving homeomorphism between S^3 does not preserve the orientation of the two knots, then we cannot say they are equivalent.

Definition 2.1.2. Two embeddings $\varphi_0, \varphi_1: K \rightarrow X$ are said *isotopic* if there is a family of embeddings $\varphi_t: K \rightarrow X$ relating them and such that $\varphi: K \times [0, 1] \rightarrow X \times [0, 1]$, $\varphi(x, t) = (\varphi_t(x), t)$ is a PL.

Two PL homeomorphisms h_0 and h_1 of X are isotopic if there is a family of homeomorphisms $h_t: X \rightarrow X$, $t \in [0, 1]$ such that $h: X \times [0, 1] \rightarrow X \times [0, 1]$, $h(x, t) = (h_t(x), t)$ is PL.

The next result allows us to introduce a simpler way to check if two links are equivalent, only by checking if a homeomorphism that takes one link into the other is isotopic to the identity.

Proposition 2.1.1. *A PL homeomorphism of S^3 preserves orientation if, and only if, it is isotopic to the identity.*

Gordon M. Fisher proved an even more general result that says that for any closed 3-manifold, the identity component of the group of all homeomorphisms $G(M)$, in other words, the largest connected subspace of $G(M)$ containing Id_M , equals the group of deformations of M , i.e., the group of all homeomorphisms of M that are isotopic to the identity. This can be found in the paper "*On the Group of All Homeomorphisms of a Manifold*".

Then we can say by everything said before that two links are equivalent if there exists an homeomorphism isotopic to the identity taking one link into the other. Sometimes we can be more general regarding the equivalence of two links by allowing the homeomorphisms revert orientation, then we say that the two links are of the same *type*. In this sense, any link can be equivalent to its mirror image.

2.2 Diagrams of Knots and Links

Let's now formally define the diagram of a knot or link. Note that these objects "live" in a three dimensional space, S^3 or $\mathbb{R}^3$, however, we can represent them in the two dimensional space S^2 or $\mathbb{R}^2$. Therefore, we shall now give a rigorous description of how to represent knots and links, this will be very useful throughout this text.

The representation of links are actually projections of such links into planes contained in $\mathbb{R}^3$. Such a projection is called a *good projection* when it is an immersion with only double point singularities. Intuitively, we have that a good projection is a closed curve in which its intersections (if it exists) occur only once in the same point.

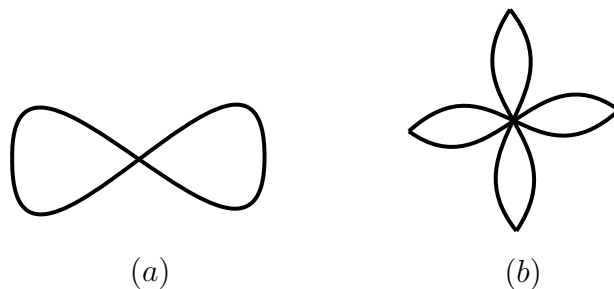


Figure 2.1: (a) is a good representation and (b) is not a good representation

Given a plane in $\mathbb{R}^3$ and a link, we can always move the link through a sequence of isotopies in such a way that we get a different representation in the plane, for the same link. And if we take a moment to think about it, the set of all planes in $\mathbb{R}^3$ that we can project the link to get a good projection is dense in the set of all planes, which can be identified with $\mathbb{R}P^3$.

Notice that in order to give a full description of the link in its representation we must add some information to the planar representation, this required information is of which part of the double point (if it exists) is passing under and which one is passing over, we do that by adding a cut in the part that passes under, and we keep a continuous curve on the part that passes over. This double points together with the under-over information is called a *crossing*. A good representation, together with the crossing information is called a *link diagram*. Figure 2.2 below, shows the diagram of the trefoil knot and the figure eight knot.

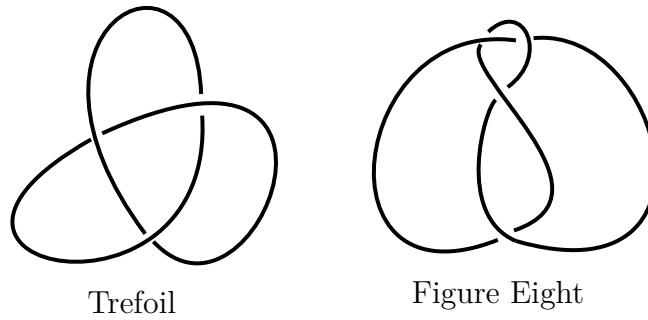


Figure 2.2: The trefoil and the figure eight.

Intuitively we can see that moving parts of the knot without passing the knot through itself gives different links, but representing the same knot. Just imagine a rope, a physical loop, exactly like a knot, of course we can move this knot around the way we want, but we are not allowed to cut this rope anywhere, the shadow it makes on a surface is like the diagram we defined above, so it is easy to see how different movements on the knot affects the diagrams. There are three main moves on a knot connecting equivalent diagrams, these moves are the so called *Reidemeister moves*. The following theorem describes how we can connect equivalent diagrams, i.e., diagrams of the same knot by means of these three fundamental moves.

Theorem 2.2.1. *Two diagrams of the same link can be obtained from one another by a sequence of isotopies and movements of the following kind:*

1. R_1 : 
2. R_2 : 
3. R_3 : 

these movements are called the Reidemeister moves.

Proof. Lets consider two link diagrams that are isotopic to each other. The isotopy between them can be decomposed into a sequence of small *triangular moves* when taking small enough subdivisions of the link strands.

These triangular moves are constructed in the following way: consider a point in the complement of the link (with a PL structure). Now, take the polygonal version of the link, in such a way that the triangle formed by two consecutive vertices of the link and the point in its complement does not intersect any other part of the link. Thus, we can substitute this side by the two other sides of the triangle, Figure 2.3 shows exactly how to do that:

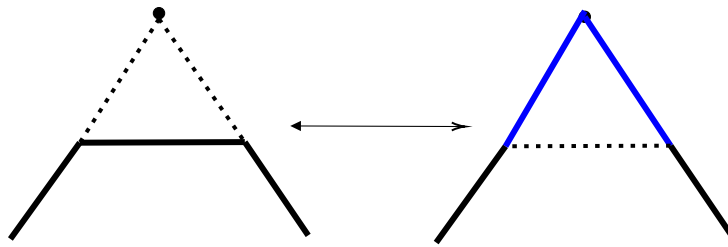


Figure 2.3: Triangular move

As the arrow indicates, we can also do the inverse procedure.

Now, taking $T \subset \mathbb{R}^3$ a triangle, let abc be its vertices in the planar representation, where b and c are consecutive vertices. If $c \in ab$, i.e., if the triangle is degenerate, then the projections of T does not change under triangular moves. Then we say abc is a *non-trivial triangle*.

Now, let's divide the triangle abc into smaller triangles. And let's consider the case where abc has the minimum amount of crossings with the rest of the diagram, we have the following cases:

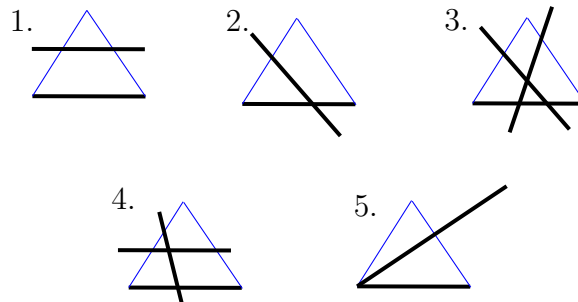


Figure 2.4: 5 different ways of triangular moves, with minimal number of intersections.

Note that any other side on the diagram has one of the 5 types of crossings as in Figure 2.4. Now, observe that each one of these moves correspond to the Reidemeister moves, in fact:

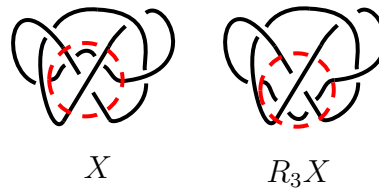
1. corresponds to R_2 ;
2. is an isotopy;
3. corresponds to R_3 ;
4. is a composition of R_3 with an isotopy;
5. corresponds to R_1 .

To see that, just notice that the blue segments are going to be the new sides of the diagram after the triangular moves, while the black segments are the old sides of the diagram. $\square$

We introduce now some algebraic notation for the Reidemeister moves, so that we can work with them without making use of diagrams by doing computations. Note that for each Reidemeister move we have two forms, if we get back to Theorem 2.2.1 we see that each move has a "before and after" version of the move, so let's name them as respectively:

1. Reidemeister move R_1 : $\ell \leftrightarrow R_1\ell$;
2. Reidemeister move R_2 : $u \leftrightarrow R_2u$;
3. Reidemeister move R_3 : $X \leftrightarrow R_3X$.

A useful remark we should make here is that when talking about the Reidemeister moves by using either diagrams or the notations introduced above, we must always remember that we are not talking about the moves as being complete diagrams, actually they are part of a diagram, remember that in the definition of a knot (or a link) we have that a link is an embedding of S^1 and therefore, the representation used on the Reidemeister moves is clearly not a full representation of the diagram. Instead, it's useful to think about the Reidemeister moves as local movements on a full diagram, and hence if we take for example, X and R_3X , we are referring to two diagrams that are identical except locally around the part of the first diagram that contains X while, in this same part, the second diagram, contains R_3X . The following picture represent exactly that:



We must also note that there are other types of Reidemeister moves, these other types are sometimes composition of the ones we've established, while others are just equivalent ways of representing the same moves, hence. The next figure gives some examples of these other kinds of Reidemeister moves.

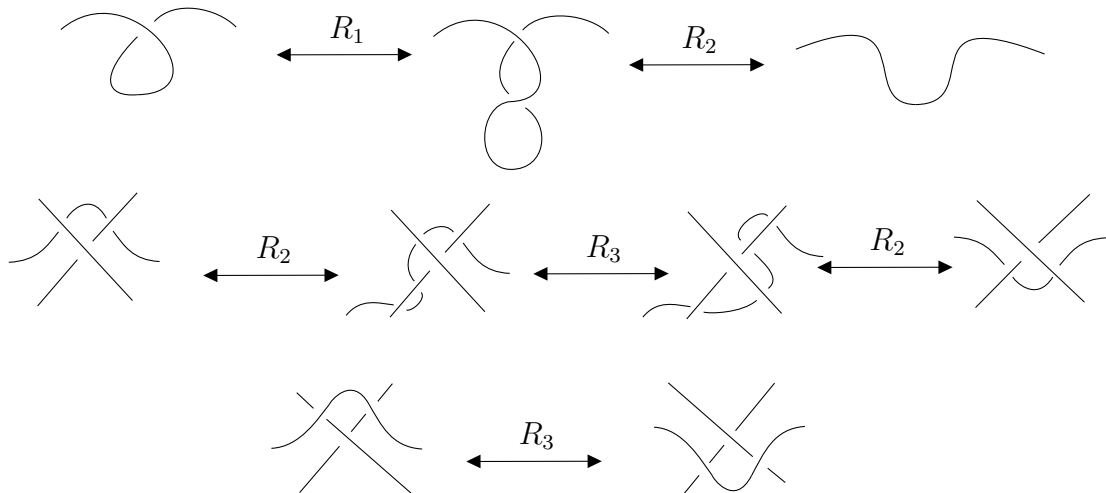


Figure 2.5: Different kinds of Reidemeister moves.

With the discovery of Theorem 2.2.1 an interesting question arise and this question was answered by J.W. Alexander and B.G. Briggs, which was concerned in discover the maximum number of Reidemeister moves necessary to connect a diagram of the unknot with n crossings into the known 0 crossings diagram. Some work showed that the upper

limit to this number is of order 2^{cn} with $c < 10^{11}$ which is a ridiculously big number. Note that given a number n we can only represent the unknot in finite number of diagrams with n crossings, thus, the 2^{cn} work, gives us an algorithm to finding out if a diagram is an unknot diagram.

The following diagram is an unknot diagram.



Figure 2.6: An unknot diagram

Goeretz showed that in order to get to the zero crossing diagram, the diagram on Figure 2.6 must, first, pass through a diagram with more than n crossings. It is still an open problem to know for which kind of diagram we must pass through before connecting two equivalent diagrams by means of Reidemeister moves. However, the problem of finding out when two diagrams are equivalent is solved, but unfortunately it is not so effective.

2.3 Knot Complements

It is common to work with the complement of knots, for example, to compute what is called the *knot group*, we need to calculate the fundamental group of the complement of a knot, for example, if we take a knot K embedded in S^3 , then the knot group of K will be the fundamental group $\pi_1(S^3 - N(K))$ of the complement $S^3 - N(K)$ of K . Here $N(K)$ denotes the *tubular neighborhood* of the knot K . An interesting fact is that, since the fundamental group of a topological space is a topological invariant, so is the knot group. Of course we can also study the groups of links in the same fashion.

Note that if two knots have homeomorphic complements then they have the same fundamental group, this is due to the fact that the knot group is the fundamental group of the knot complement and the fact that homeomorphic spaces have the same fundamental groups.

One way we can compute the knot group is through the *Wirtinger Presentation*. We now describe how to obtain such presentation from a knot (or link) diagram, then we present an argument to convince ourselves that this presentation is in fact a presentation for the knot group.

We start by naming each arc on the knot diagram as a_i , and we do this so the consecutive arcs that are separated by a crossing have consecutive names, i.e., a_i and a_{i+1} , for example. Now we fix a point P on the complement of the knot (which we shall denote as $S^3 - N(K)$ as before), and for each arc a_i we make a loop based on P called α_i for the arc a_i .

We can think of these loops as meridians of the arcs, since we can simply take a meridian for each arc and then connect it to P via a line that avoids the tubular neighborhood

of the knot. We also ask that these loops are oriented in a way that the crossing generated by passing α_i under a_i is a positive crossing. We may do that, intuitively, by using the right-hand rule, i.e., we point our thumb of our right hand towards the direction of a_i given by the orientation of the knot, and the orientation we must give to α_i is such that it goes in the direction the rest of our fingers.

Take now a negative crossing in the given knot, call a_j the over strand and α_i and α_{i+1} the two arcs corresponding to the under strand, we can easily see that the word $\alpha_j \alpha_i \alpha_j^{-1} \alpha_{i+1}^{-1}$ can be entirely deformed to the base point P on the complement of K , since it is a loop that goes under the whole crossing, i.e., it passes under both the over strand and the under strand, hence, it is the trivial loop, i.e.:

$$\alpha_j \alpha_i \alpha_j^{-1} \alpha_{i+1}^{-1} = e \Leftrightarrow \alpha_j \alpha_i = \alpha_{i+1} \alpha_j$$

this is a set of relations that the group generated by the meridians must have, for each crossing.

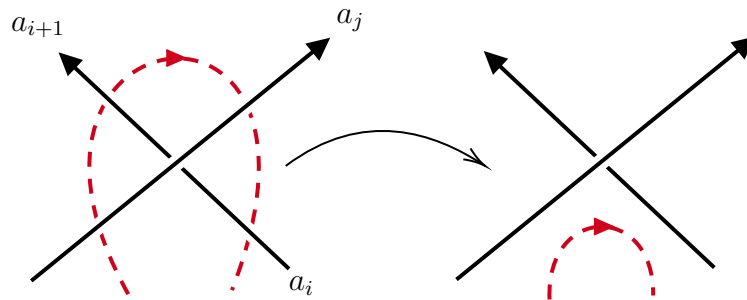
The other set of relations is given in a similar fashion, but now looking at a positive crossing of the knot projection, and it is given by:

$$\alpha_j \alpha_{i+1} = \alpha_i \alpha_j.$$

Thus, the Wirtinger presentation, is the group generated by the α_i 's and subjected to these two set of relations given at each crossing of the knot. Of course, this whole construction can be extended to the context of links.

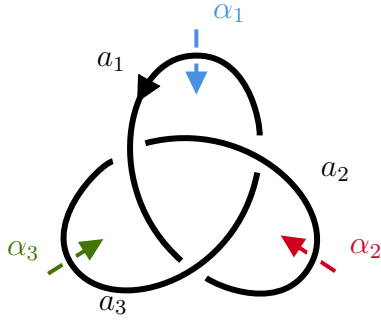
Now, we need to see that this is in fact a presentation for the knot group. To see this, note that if we take any loop based at P on the complement of K can be broken into a composition of the α_i 's by doing the following: every time the loop goes under some arc of the knot, we deform it back to P and then back again to where it was before the deformation, doing this, we are deforming the whole loop into a composition of loops that go around only one arc, which are precisely the α_i loops, hence, every loop based on P is a word on the generators of the Wirtinger presentation.

To see that the relations hold in the fundamental group of the knot complement, we just repeat the argument that if a loop starts at P , goes entirely under a crossing, i.e., it passes under both the over and under strands of the crossing, then it wraps around the crossing and go back to P , it can be deformed into the trivial loop. Also, by using the previous break down of a loop into a word on the generators, we break this loop and note that its corresponding word is precisely the word that gives us the relations on the presentation. The following picture shows this whole process:



Example. Let's compute the Wirtinger presentation of the group of the trefoil knot. To simplify drawings we will draw the meridians as small arrows perpendicular to the arcs

they correspond to, the reader may imagine the base point on the complement of the knot as being his eyes, or a point just above the plane of projection of the knot. Let's consider the orientations and notations on the following picture:



From the relations described above applied to each crossing on the picture (note that each crossing corresponds to a negative crossing), we get the following:

$$\alpha_3\alpha_1 = \alpha_2\alpha_3, \alpha_1\alpha_2 = \alpha_3\alpha_1 \text{ and } \alpha_2\alpha_3 = \alpha_1\alpha_2$$

now, note that we can write one of the generators in terms of the other two, for example $\alpha_3 = \alpha_2^{-1}\alpha_1\alpha_2$, hence, we can write the group presentation as the group generated by two elements, α_1 and α_2 , and satisfying the relation $\alpha_1\alpha_2\alpha_1 = \alpha_2\alpha_1\alpha_2$, obtained by putting $\alpha_3 = \alpha_2^{-1}\alpha_1\alpha_2$ in any of the relations above, hence:

$$\pi(T) = \pi_1(S^3 - N(T)) = \langle \alpha_1, \alpha_2 \mid \alpha_1\alpha_2\alpha_1 = \alpha_2\alpha_1\alpha_2 \rangle.$$

Now we turn our attention to composite knots and prime knots, we will come back to prime knots later, with a more formal definition, but for now we will look at the basic idea. The main idea is to make an analogy to integer numbers, in the integers we have that a number is prime when it cannot be written as the product of other prime numbers different from the initial one, while composite numbers can, for example $2 = 2 \cdot 1$ therefore 2 is prime, while $4 = 2 \cdot 2$, therefore 4 is composite. The same idea applies to knots.

Lets define a connected sum of two knots. Suppose we have two oriented knots K_1 and K_2 that are disjoint. Now take piece of ribbon with extremities on both knots and with the orientations on the boundaries opposite to each other. Now excise small arcs from each knot and replace them with the boundary of the ribbon. This new knot obtained this way will be denoted as $K_1 \# K_2$ and called *connected sum* of K_1 and K_2 . Figure 2.7 represents this process.

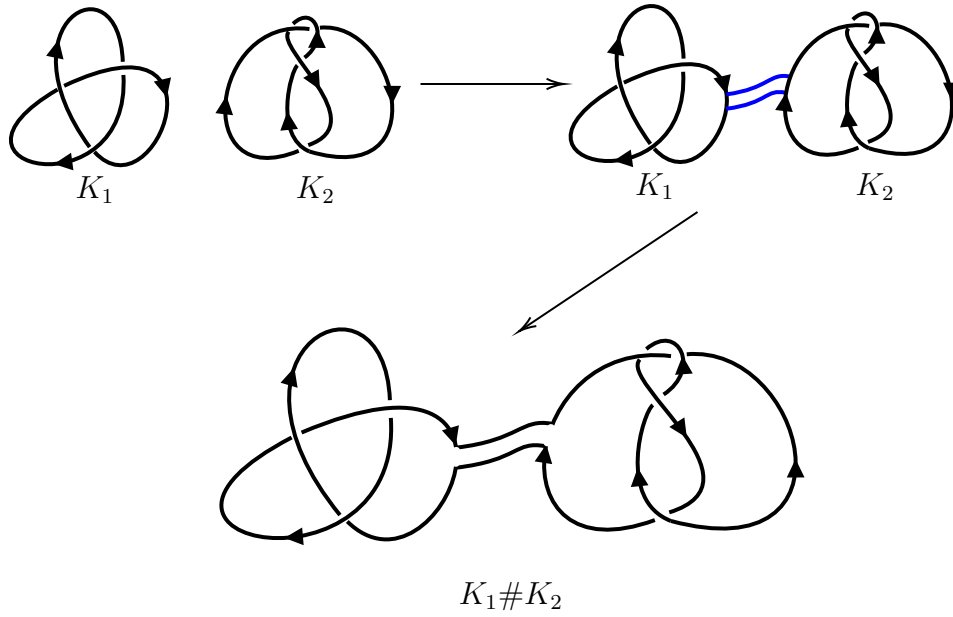


Figure 2.7: Connected sum of a trefoil (3_1) and a figure eight knot (4_1).

In an analogous way we can define primeness of a knot by saying a knot is *prime* when it is not composite in respect to the connected sum, i.e., when a knot can be written as the connected sum of two knots, then it is said to be a *composite knot*.

Theorem 2.3.1. *Two prime knots with isomorphic groups have the homeomorphic complements, such homeomorphism preserves orientation.*

Note that the hypothesis of the Theorem 2.3.1 asks the primeness of the two knots, meaning that isomorphic knot groups does not necessarily implies homeomorphic complements. However, the knot group of a given knot determines the type of each one of its components, thus taking connected sums of knots with different orientations generates examples of different knots with the same knot group.

2.4 Triangulations on Surfaces

In two dimensions, we are going to "measure" the complexity of a triangulation, by counting the number of vertices there are in a given triangulation for a surface. Remember that a triangulation can be viewed as a pair $(\mathcal{K}, \phi)$, in which $\mathcal{K}$ is a simplicial complex and ϕ is a homeomorphism between $|\mathcal{K}|$ and the surface. So we say a triangulation in two dimensions is more complex than other if it has more vertices.

In two dimensions, we have some moves to create new triangulations, these are called the *Pachner moves*. The first two Pachner moves are the *rotation* and the *barycentric subdivision*. We now, introduce briefly the barycentric subdivision.

Let s^n be an n -dimensional simplex with vertices $a_0, a_1, \dots, a_n$. The *barycenter* of s^n is the point in s^n given by

$$b(s^n) = \frac{1}{n+1} \sum_{i=1}^n a_i.$$

Now, we shall construct recursively the barycentric subdivision, which we will call Sd . First, set $Sd(s^0) = s^0$, for any s^0 . Suppose now $Sd(s^{n-1})$ is defined. Then taking t^{n-1} any $n-1$ -simplex, then $Sd(t^{n-1})$ is the collection of all $(n-1)$ -simplexes contained in t^{n-1} . Let now $\dot{s}^n$ be the collection of all $(n-1)$ -faces of s^n and define:

$$Sd(\dot{s}^n) = \bigcup_{t^{n-1} \in \dot{s}^n} Sd(t^{n-1}).$$

Thus, $Sd(s^n)$ consists of all n -simplexes of the form $(b(s^n), t_0, \dots, t_{n-1})$, where $(t_0, \dots, t_{n-1})$ is a $(n-1)$ -simplex contained in $Sd(\dot{s}^n)$. Hence we have defined the barycentric subdivision of s^n . Figure 2.8, is an example of a barycentric subdivision of a 2-simplex:

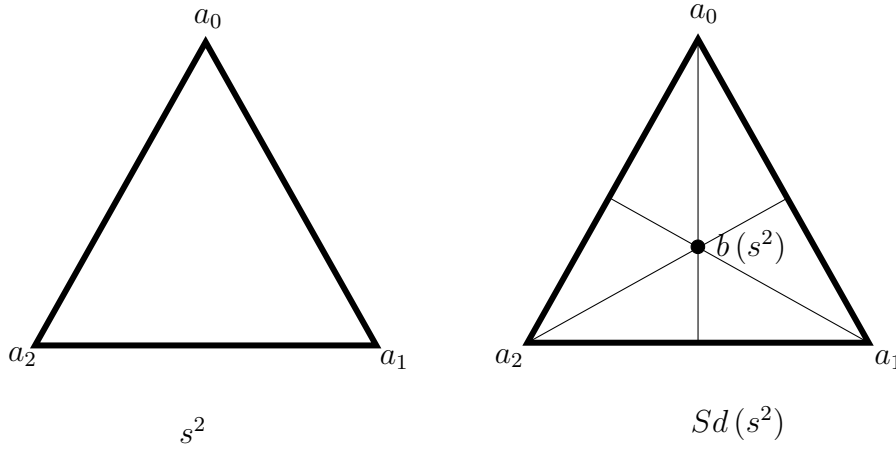


Figure 2.8: Barycentric subdivision

Now we can finally see what are the Pachner moves in triangulations. As we have already defined, they are a rotation and a barycentric subdivision. The barycentric subdivision has already been described, and visually stated, now the rotation is defined as follows: take two adjacent 2-simplexes, it forms a parallelogram now take the common diagonal that are the intersecting faces of the two simplexes, and substitute it by the other diagonal. Figure 2.9 shows exactly that:

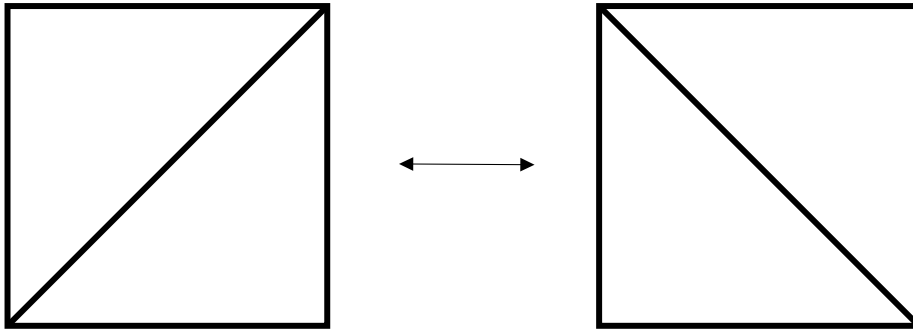


Figure 2.9: Rotation

Its important to notice that it is not always true that two triangulations of some given closed surface, both with the same number of vertices are equivalent by means of rotations, meaning that we can get one from the other only by rotations, unless we allow some kind of degenerate rotations.

2.5 The Linking Number

Taking any oriented diagram, different from the trivial one, we always have some number of crossings. These crossings are of two different types as in Figure 2.10 below:



Figure 2.10: Types of crossings

Now imagine taking the over strand of the crossing and rotating it by $\frac{\pi}{2}$ radians anticlockwise, if the orientation corresponds, or it is parallel, to the orientation of the under strand, we set the *sign* of this crossing to be positive, and we denote it by $\varepsilon(L_1) = +1$. In the same way, if the orientation of the over strand is opposite to the orientation of the under strand, we set the sign of this crossing to be negative, and we denote it by $\varepsilon(L_0) = -1$.

Generally, the literature on this topic uses notations in diagrams, however we are going to use only algebraic notations and use diagrams only to represent some ideas.

We introduce now one of the most important concepts on the study of knots, the concept of *linking number*. This is a tool that relates an oriented link diagram with to a number, that can be understood as an amount of times one strand of a link winds around the other strand of the same link, so we can conjecture that if a link diagram is a diagram for the trivial link, then its linking number should be zero, but unfortunately that is not what happens, and we will get to that.

Suppose a diagram of a link with two components L and K , we denote by $\overrightarrow{D_L}$ and $\overrightarrow{D_K}$ the oriented diagrams of L and K respectively. Now, denoting by $\overrightarrow{D_L} \cap \overrightarrow{D_K}$ the set of all crossings between L and K , then we define the linking number between L and K as:

$$lk(L, K) = \frac{1}{2} \sum_{p \in \overrightarrow{D_L} \cap \overrightarrow{D_K}} \varepsilon(p).$$

The importance of the linking number is clear on the next theorem, since it shows that the linking number is a topological invariant.

Theorem 2.5.1. *Any two equivalent diagrams have the same linking number.*

Proof. Since two equivalent diagrams can be connected by means of Reidemeister moves, it is sufficient to check that the linking number is invariant by the Reidemeister moves. We do that:

1. The R_1 move. Since the first Reidemeister move occurs in only one strand, the linking number is zero before and after the move, remember that by definition, the linking number is computed only in crossings of strands originated from different components, the crossings that occur on the same component has sign zero.

2. The R_2 move. Consider the following diagram with the corresponding notations:

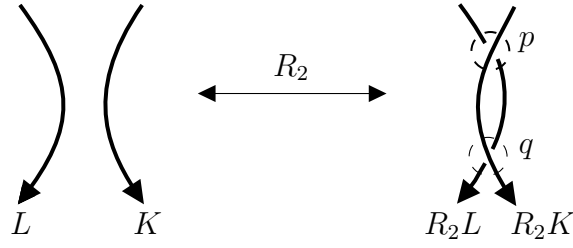


Figure 2.11: Linking number of R_2 move

Now, note that the crossings p and q will correspond respectively to a crossing of the type L_1 and L_0 . Then by definition, we have that $lk(L, K) = 0$ since L and K are disjoint, and $lk(R_2L, R_2K) = \frac{1 + (-1)}{2} = 2$, then in this configuration, the linking number is invariant by Reidemeister move R_2 . For any other configuration, the proof is analogous.

3. The R_3 move. We now color the strands on the diagram to represent which strand originates by which component. Since the linking number is defined for two components and the R_3 move occurs in three strands, two of them will be colored by the same color, meaning that they come from the same component of the link. Again, consider the notation on the following diagram:

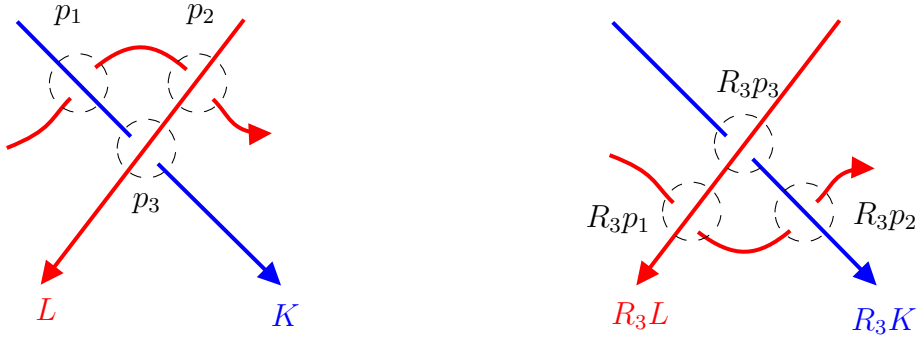


Figure 2.12: Linking number of R_3 move

Now, observe that the crossings p_2 and R_3p_1 are crossings in the same component of the link, L and R_3L respectively, therefore their signs are zero. Now, the crossings p_1 , p_3 , R_3p_2 and R_3p_3 are all of the same type, L_1 , then we have:

$$lk(L, K) = \frac{+1 + 1}{2} = 1 \quad \text{and} \quad lk(R_3L, R_3K) = \frac{+1 + 1}{2} = 2$$

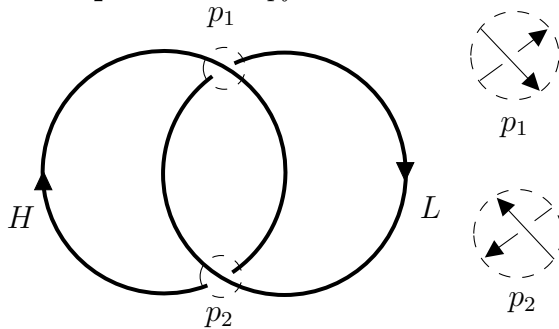
Hence, $lk(L, K) = lk(R_3L, R_3K)$ and the linking number is invariant by Reidemeister move R_3 . Any other configuration of the diagram is analogous to this one.

So, since the linking number is an invariant of Reidemeister moves, we conclude the result. $\square$

Although it looks like all diagrams with zero linking number are diagrams of the trivial link (a link with two components, each being the unknot), that does not happen. The fact that the crossings of the same component do not affect the linking number, allows us to create a link such that the two components, in terms of linking number, are not "actually linked", i.e., they have zero linking number, but they cannot be untangled, this is what happens with the Whitehead's link. Another way to see that, is by taking now a link with three components, such that, as seen as a three component link, it is "linked" in the sense that we cannot separate the components, but when we look to two of its components, they are, in fact, unlinked, so their linking number will be zero. This happens because the linking number is defined only for two components.

The following examples show exactly what we have said about links that are not unlinked but have zero linking number.

Example. The *Hopf link*:

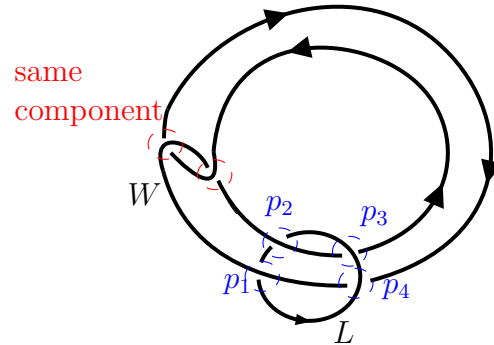


Notice that the crossings p_1 and p_2 are both of the type L_1 , thus we have that $lk(H, L) = \frac{1}{2}(\varepsilon(p_1) + \varepsilon(p_2)) = \frac{1}{2}(1 + 1) = 1$, hence, the Hopf link has linking number 1

Example. The *Whitehead's link*:

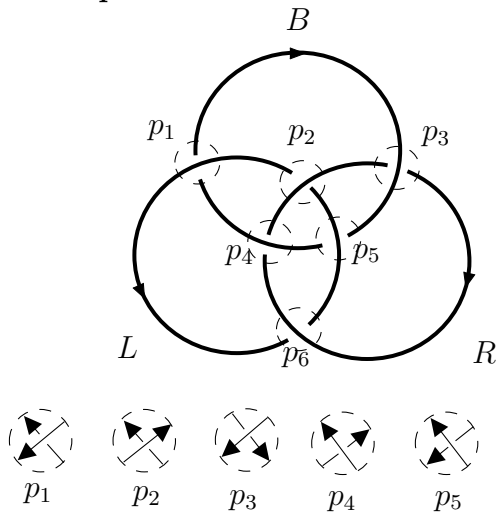
Now, we have that the crossings p_1 and p_4 are of the type L_1 , while p_2 and p_3 are of the type L_0 , therefore, we have

$$lk(W, L) = \frac{1}{2}[\varepsilon(p_1) + \varepsilon(p_2) + \varepsilon(p_3) + \varepsilon(p_4)] = \frac{+1 - 1 - 1 + 1}{2} = 0.$$



Thus the Whitehead's link is a link with zero linking number, but it is not a diagram for the trivial link with two components, as we have noticed before.

Example. The *Borromean link*:



Since now we have three components, we must check the linking number for each two of them. Note first that the crossings p_1 , p_4 and p_6 are of the type L_0 while the crossings p_2 , p_3 and p_5 are of the type L_1 . Thus we have the following, for the B and R components:

$$lk(B, R) = \frac{\varepsilon(p_3) + \varepsilon(p_4)}{2} = \frac{1 - 1}{2} = 0$$

now, for the components B and L :

$$lk(B, L) = \frac{\varepsilon(p_1) + \varepsilon(p_5)}{2} = \frac{-1 + 1}{2} = 0$$

and finally for the last two components, R and L :

$$lk(R, L) = \frac{\varepsilon(p_2) + \varepsilon(p_6)}{2} = \frac{1 - 1}{2} = 0$$

Hence, although the Borromean link is not the trivial link with three components, every two component is not linked to each other, i.e., if we delete any one of the three components, we get the trivial link with two components. This happens because if we take any component, it is not linked to itself, unlike the Whitehead's link.

We will now note some interesting properties of the linking number. The first one is the fact that the linking number is always an integer, independently of how many crossings a diagram have, the linking number is always going to be an integer, this can be easily seen by noting that if we allow a crossing to change sign, for example a positive sign to a negative one, then the amount of positive crossings, say p are going to reduce one unity, while the amount of negative crossings, say q are going to grow by one unity, so in the overall picture we have, supposing a link with two components L and K , with L' and K' being the link with the changed crossings:

$$lk(L', K') = \frac{1}{2}((p - 1) - (q + 1)) = \frac{1}{2}(\sum(p - q) - 2) = lk(L, K) - 1$$

and since the linking number of the trivial link is zero, then allowing a link to become the trivial link only by changing crossings, then the total amount of the original linking number must be an integer.

A nice consequence of this (which is also a way to see why this holds) is that we can compute the linking number only by looking at the negative or the positive crossings, then we define

$$lk'(L, K) = \sum_{p \in [D_L \cap D_K]_{L > K}} \varepsilon(p)$$

where $[D_L \cap D_K]_{L > K}$ represents the set of all crossings between the diagrams D_L of L and D_K of K , where L passes over K , and $lk'(L, K) = lk(L, K)$ (a similar construction can be made for ne crossings of L passing under K). And in fact, as we have seen, changing a crossing from a positive to a negative changes the linking number in -1 while changing a crossing from a negative one to a positive one, changes the linking number in 1 (this follows in the same fashion as the first one), hence, either by looking only to positive (or negative) crossings, or to both crossings at the same time, by changing all the positive (or negative) crossings must give us the linking number of the trivial link, i.e., zero. So the numbers $lk(L, K)$ and $lk'(L, K)$ must be the same.

2.6 Seifert Surfaces

With the intent of giving a topological interpretation to the linking number, we define a certain kind of two dimensional surface embedded in S^3 . These are a very special kind of surfaces that are close related to knots and links. Suppose an oriented surface which its boundary is a knot, for example the trefoil knot, or even a link, like the Hopf link, these surfaces receive the name of *Seifert surfaces* due to the German mathematician Herbert Seifert.

It may be a little messy to imagine such a surface, but we can naively try to define it through the knot diagrams. Suppose we color the diagram of a link in a chessboard fashion, like Figure 2.13. Now, imagine each colored part of this diagram is part of the surface and each crossing corresponds to a kind of twist in the surface. This is a good first look in how to create such Seifert surfaces, but unfortunately, every time we do that, we divide S^3 into two parts and one of them contains the point at infinity (seen as the one point compactification of $\mathbb{R}^3$), and the constructed surface may not be orientable, and this is a problem.

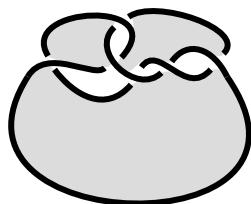
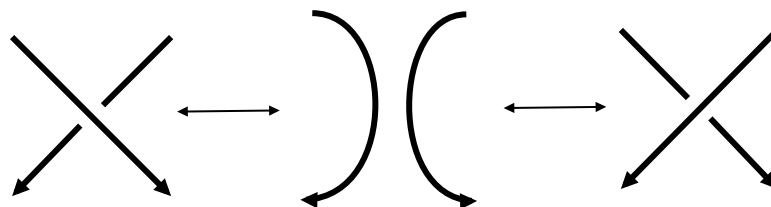


Figure 2.13: A first try on constructing a Seifert surface.

The following theorem not only gives us an affirmative answer to the question: given a knot or link, will there always be an associated Seifert surface? But it also gives us a way to construct such surfaces, and the good part of it is that, since we are constructing a Seifert surface, it will definitely be orientable, unlike our first try on solving this problem.

Theorem 2.6.1. *Any link in S^3 has an associated Seifert surface.*

Proof. Given a link L and a oriented diagram $\vec{D}$ for L , take each crossing on the link diagram and smooth the crossing in a way that the orientation of the smoothed crossing is compatible with the orientation of the rest of the diagram, as follows:



By doing this with every crossing, we get a diagram with a finite number of disjoint copies of S^1 , these are called *Seifert circles*. Consider the Seifert circles as boundaries of two disks D^2 in different heights on $\mathbb{R}^3$, i.e., each Seifert circle bounds disjoint disks D^2 in $\mathbb{R}^3$. Now, where the crossings were, we connect between each Seifert circle a strip twisted in π radians and respecting the position of the crossing, meaning that the part of the strip passing over the other corresponds to the strand of the crossing passing over the other strand, and vice-versa.

Finally, the obtained surface is orientable by construction and its boundary is, also by construction, the given link. Hence we have constructed a Seifert surface. $\square$

Let's see now an example of this construction with the Figure Eight knot, so we can visually understand what is happening.

Example. The Seifert surface of the Figure Eight knot is constructed as in Figure 2.14, notice the Seifert circles after the crossing smoothing, and also the twisted stripes where the crossings were before the smoothing.

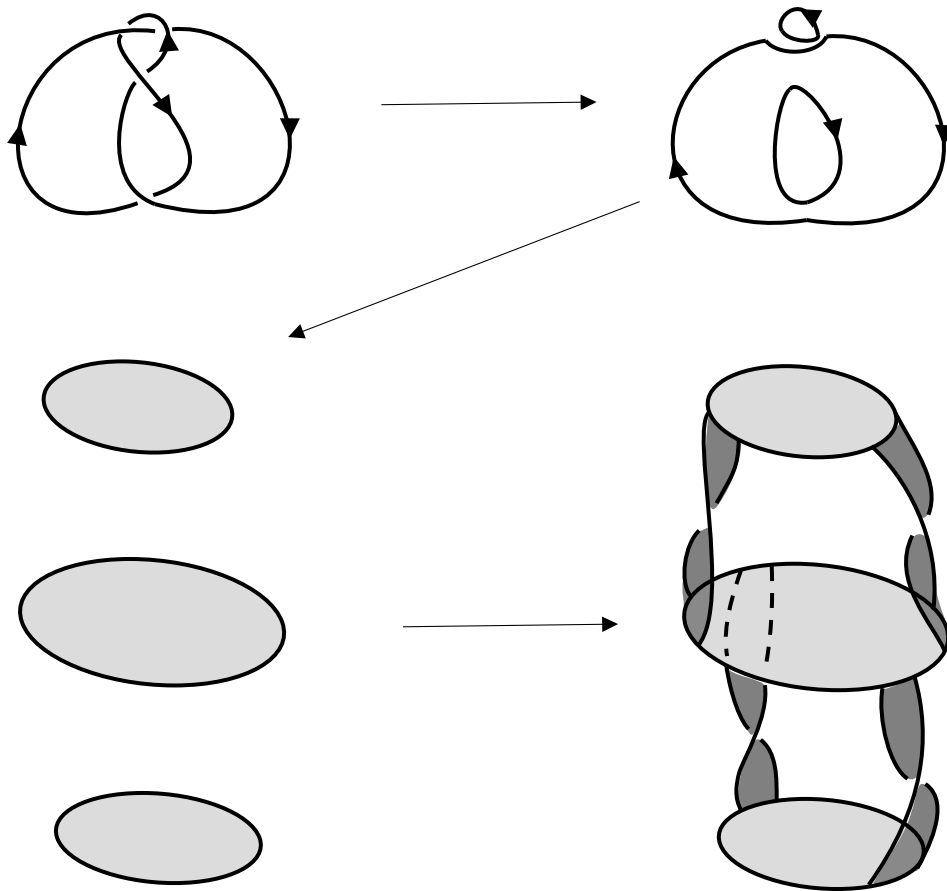


Figure 2.14: The Seifert surface of the Figure Eight knot.

Example. Another example, a little more simple is the Seifert surface of the trefoil, again, the same construction introduced on the proof of Theorem 2.6.1 is done here:

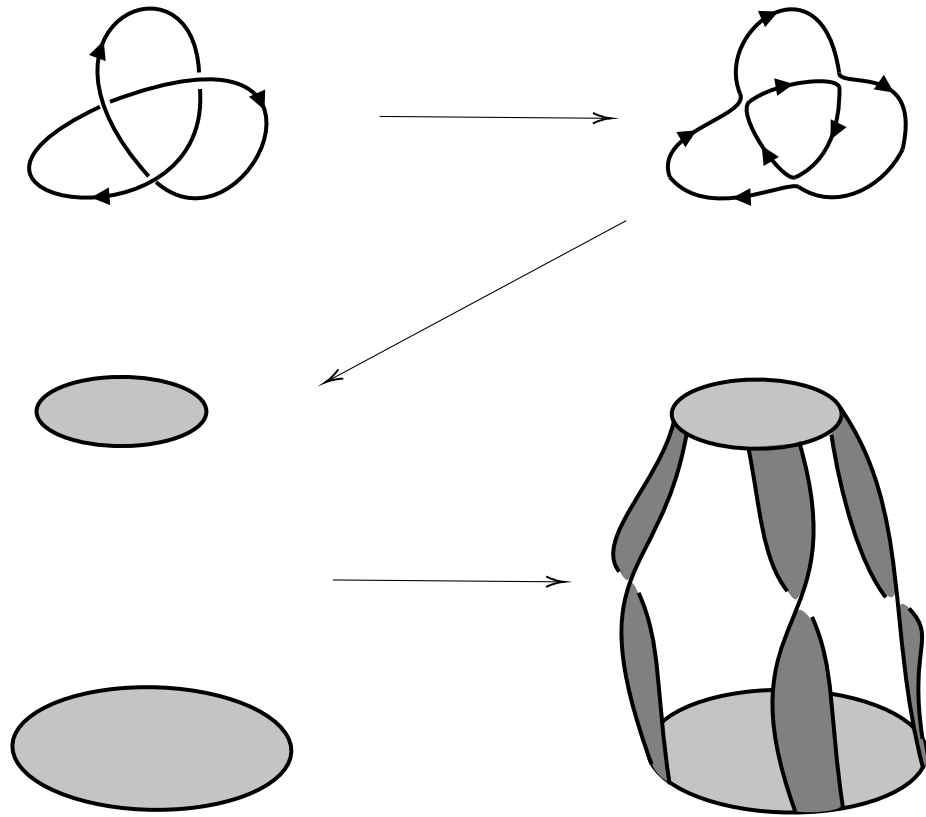


Figure 2.15: The Seifert surface of the Trefoil knot.

Now, we can introduce a very important theorem that gives us some topological, and even physical interpretation of the linking number.

- Theorem 2.6.2.** 1. Let K be an oriented knot in S^3 . Then there exists a isomorphism $H_1(S^3 - N(K)) \approx \mathbb{Z}$ with generator given by the homology class of a simple loop μ on the boundary, and which is the boundary of a two disk on the cylinder $N(K)$ and intersects K at only one point (μ is called a meridian). Thus the class of all simple oriented curves $L \subset S^3 - N(K)$ in $H_1(S^3 - N(K))$ is $lk(K, L)lk(K, \mu) \in \mathbb{Z}$.
2. Let F be an oriented surface in S^3 which boundary is the knot L . Any other knot K can be moved to become transversal to F . For each intersection $q \in K \cap F$ we associate a number $\tilde{\varepsilon}(q)$ as follows: the orientation of F together with the orientation of K gives an orientation for $\mathbb{R}^3$ at q , if this orientation is compatible with the usual orientation of $\mathbb{R}^3$, then $\tilde{\varepsilon}(q) = +1$, if it is not compatible, then $\tilde{\varepsilon}(q) = -1$. Then:

$$lk(K, L) = \sum_{q \in K \cap F} \tilde{\varepsilon}(q).$$

3. Let $\varphi_1, \varphi_2: S^1 \rightarrow \mathbb{R}^3$ be parameterizations of L and K , we have that $\Gamma(s, t) = \frac{\varphi_1(t) - \varphi_2(s)}{\|\varphi_1(t) - \varphi_2(s)\|}$ is the Gauss map from $S^1 \times S^1$ to S^2 . Then

$$lk(K, L) = \deg \varphi = \frac{1}{4\pi} \int_{K \times L} du \left(\frac{dw \times (u - w)}{\|u - w\|^3} \right)$$

Proof. 1. Lets first prove 2. Consider the following definition for the linking number:

$$lk'(L, K) = \sum_{p \in [\vec{D}_L \cap \vec{D}_K]_{K > L}} \varepsilon(p)$$

where $[\vec{D}_L \cap \vec{D}_K]_{K > L}$ represents the set of crossings of K and L where K passes over L . Then since this is another way of writing the linking number, it is of course a topological invariant.

Consider now the Seifert surface S of L associated with the diagram $\vec{D}_L$ constructed as in the Theorem 2.6.1. Now, consider the Sifert circles in the construction, we can, by means of isotopies put K over S except near a small neighborhood of each crossing. The number associated to L passing under K coincides with $\varepsilon(p)$, thus we have proven this claim.

2. Consider the cover $\{N(K), S^3 - N(K)\}$ for S^3 , then, the Mayer-Vietoris sequence for this cover is

$$\cdots \rightarrow H_2(S^3) \rightarrow H_1(\partial N(K)) \rightarrow H_1(S^3 - N(K)) \oplus H_1(N(K)) \rightarrow H_1(S^3) \rightarrow \cdots$$

Since $H_2(S^3) = H_1(S^3) = \{0\}$, it follows from the above sequence that $H_1(\partial N(K)) \approx H_1(S^3) \oplus H_1(N(K))$. But, note that $\partial N(K) \approx \mathbb{T}^2$ and that $H_1(\mathbb{T}^2) \approx \mathbb{Z} \oplus \mathbb{Z}$. Now, since by the definition of the Mayer-Vietoris sequence, the isomorphism is induced by the inclusion map of $\partial N(K)$ into $N(K)$ and, note that a meridian in $\partial N(K)$ is zero in the homology group of $N(K)$, then it follows that $H_1(S^3 - N(K)) \approx \mathbb{Z}$. Similarly, we have that $H_2(S^3 - N(K)) = H_2(S^3 - N(K)) = \{0\}$.

We now have that $\mathbb{Z} \approx H_1(S^3 - N(K)) \xrightarrow{\psi} H_0(S^3 - N(K)) \approx \mathbb{Z}$ can be constructed as follows: take a 1-cycle $x \in S^3 - N(K)$, we must have that $[x] = 0 \in H_1(S^3)$, then there exists a 2-chain y in S^3 such that $\partial y = x$. Hence, the intersection of y with K is a 0-cycle of $S^3 - N(K)$, then the map $x \mapsto [y \cap K]$ gives us the map ψ . Note also that taking μ as in the theorem, we have $\psi(\mu) = lk(K, \mu) \in \{1, -1\}$ (note that it looks like a Hopf link locally), which is a generator for $\mathbb{Z}$ since it is a generator for $H_1(S^3 - N(K))$.

On the other hand, if L is a simple curve on the complement of K , take the Seifert surface S of L for some diagram of L , then the image of L through ψ is going to be $[S \cap K] = lk(K, L) \in \mathbb{Z}$, since the signs associated with each intersection point are the same associated with the intersection of the cycles. Thus the homology class of L in $H_1(S^3 - N(K))$ must be $lk(K, L)lk(K, \mu) \in \mathbb{Z}$.

3. The Gauss map is a continuous function that associates for each point in $S^1 \times S^1$ a unit vector that is normal in S^2 . The function $\Gamma(s, t) = \frac{\varphi_1(s) - \varphi_2(t)}{\|\varphi_1(s) - \varphi_2(t)\|}$ is a Gauss map.

Now, note that if we take $z \in S^2$ directly above (or seen as a vector, perpendicular to) the projection plane of the link, then the vector $\varphi_1(s) = \varphi_2(t)$ have the same direction as z if, and only if, $K > L$. In other words, $\Gamma^{-1}(z) = [\vec{D}_K \cap \vec{D}_K]_{K > L}$, and writing $lk(L, K) = \sum \varepsilon(p)$ in the usual way, we have that $lk(L, K)$ can be written as

$$lk(L, K) = \sum_{p \in \Gamma^{-1}(z)} \tilde{\varepsilon}(p) = \deg(\Gamma).$$

The fact that

$$\deg(\Gamma) = \int_{L \times K} du \left(\frac{dw \times (u - w)}{\|u - w\|^3} \right)$$

follows from [20] exercise number 5-32, on page 132.

□

2.7 Milnor's Invariants or The Higher Linking Numbers

Throughout this chapter we have studied about some useful concepts that will be used in the construction of link invariants in the future chapters. In this moment we will introduce yet another tool that will not be as used as the previous ones in this dissertation, but it is interesting to know that it exists.

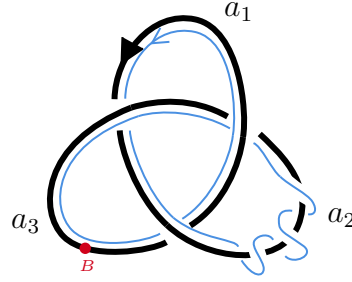
When studying the linking number, we figured out it could be understood as a measurement of how one component of a link is tangled with another component. However, we've seen links that are obviously tangled, although their linking number between any two components is zero, this is the case of the Borromean rings. This fact happens since the linking number detects entanglement between any two components, and the word "two" here is important.

It is true that any two components of the Borromean rings link are not tangled, though the three of their components are. It would be great to have a tool that detects entanglement between more than two components, that is the work we do now. We define the Milnor's invariants, which can be, as in the case of the linking number, seen as a measurement of how three or more components of a link are tangled. Of course, at the end of this section, we expect to find that this invariant when computed on the Borromean rings link gives us a value different from zero for the its three components.

First we need to define the *longitude* of a link (or a knot). Also, let's note that sometimes, for the sake of simplicity, we will restrict ourselves to the study of knots, although it will be trivially extended to the study of links in general. The longitude of a knot is a curve in the exterior of the knot that has linking number zero with the knot. It can be thought of as copy of the knot shifted a little bit far away from the knot so that the knot does not intersects this copy.

Of course, since the longitude of the knot is a loop in the complement of the knot, it can be seen as an element of the knot group, hence, it can be written as a word on $\pi(K)$. To realize the longitude as a word we follow the following steps: begin with the trivial word and a base point on the knot K , follow the strand of the knot in the direction induced by the orientation of the knot, every time we reach an undercrossing, we multiply on the right the word we had previously by the meridian corresponding to the overpassing strand if the crossing is positive, or by the inverse of such meridian if the crossing is negative.

Example. Consider T the trefoil knot oriented, with meridians and base point B , as denoted in the picture:



We may intuitively think about the curve that is just the copy of T shifted a little far away from T , in this case, we can compute the (possible) longitude¹ $\lambda = \alpha_2 \alpha_1 \alpha_3$. But, when we compute the linking number of λ with T , we see that $lk(\lambda, T) = 3 \neq 0$, hence, this λ , although a candidate for, is not a longitude. But we can easily correct this by adding 3 negative crossings, by wrapping the curve λ around some arc of T three times, as in the picture. In this case, we get the following word which correctly describes the longitude of T : $\lambda = \alpha_2 \alpha_3^{-3} \alpha_1 \alpha_3$.

Now, let L be a link with n arcs that correspond to the n meridians that generate the knot group of L , we define a function Φ defined on the set of generators of the knot group and taking values on the ring of power series $\mathbb{Z}[t_1, \dots, t_n]$ where the t_i 's are non commutative variables, in the following way:

$$\Phi(\alpha_i) = 1 + t_i$$

$$\Phi(\alpha_i^{-1}) = \sum_{k=0}^{\infty} (-1)^k t_i^k.$$

This function is called the *Magnus expansion*. Our interest here will be in the Magnus expansion of each longitude of the link. Recall that a longitude is a word in the meridians of the knot, hence, each longitude can have a Magnus expansion in the trivial way, suppose, for example, that $\lambda = \prod_{j=1}^n \alpha_j^{\nu_j}$ is a word that represents a longitude of a link, then:

$$\Phi(\lambda) = \Phi\left(\prod_{j=1}^n \alpha_j^{\nu_j}\right) = \prod_{j=1}^n \Phi(\alpha_j^{\nu_j}).$$

Let's denote as $\mu_L(i_1 i_2 \dots i_k, j)$ the coefficient of the term $t_{i_1} t_{i_2} \dots t_{i_k}$ in the Magnus expansion of the j -th longitude of the link L , consider I a list of indexes then we define the value $\Delta_L(I)$ to be the greatest common divisor of all $\mu_L(I')$ where I' is a list obtained from I by removing one index and permuting the other in a cyclic fashion, as an example, consider $I = \{1, 2, 3\}$, then

$$\Delta_L(1, 2, 3) = \gcd\{\mu_L(1, 2), \mu_L(2, 1), \mu_L(1, 3), \mu_L(3, 1), \mu_L(2, 3), \mu_L(3, 2)\}.$$

Finally, we define the *Milnor's Residue Class*, also known as *Higher Linking Number* as:

$$\bar{\mu}_L(I) \equiv \mu_L(I) \pmod{\Delta_L(I)}.$$

¹To compute this, just ignore the "zig-zag" on the third arc in the picture and follow the algorithm described in the previous text.

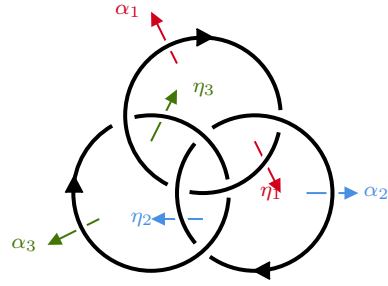
Of course, this is a general definition that works for any list I , however our main interest on the Milnor's invariants will be when I is a list of three elements, and intuitively this number corresponds to something like a triple linking number, i.e., it is a way to measure how three different components are tangled.

We now go back to the example of the Borromean rings and try to determine whether this new invariant, the triple linking number, gives us the expected result of telling us that the three components of the link are tangled while any two of them are not. But first we shall remark that, if I is a list of two elements, then each $\mu_L(I')$ is zero (it is a convention that $\mu_L(i)$ is always zero for any index i), hence $\bar{\mu}_L(I) = \mu_L(I)$, and, in this case, $\mu_L(I)$ is just the linking number between the two components list in I .

Example. Consider the Borromean rings link B and the notations represented in the following picture:

We have, of course, three longitudes, which can be determined by following the usual steps introduced earlier, thus we get:

$$\begin{cases} \lambda_1 = \alpha_2^{-1} \eta_2 \\ \lambda_2 = \alpha_3^{-1} \eta_3 \\ \lambda_3 = \alpha_1^{-1} \eta_1 \end{cases}.$$



Now, note that we can write the Wirtinger presentation of the knot group of B as:

$$\pi_1(S^3 - N(B)) = \left\langle \begin{array}{c|c} \alpha_1, \alpha_2, \alpha_3 & \alpha_1 \alpha_3 = \eta_3 \alpha_1; \alpha_2 \alpha_1 = \eta_1 \alpha_2; \alpha_3 \alpha_2 = \eta_2 \alpha_3 \\ \hline \eta_1, \eta_2, \eta_3 & \eta_2 \alpha_1 = \eta_1 \eta_2; \eta_3 \alpha_2 = \eta_2 \eta_3; \eta_1 \alpha_3 = \eta_3 \eta_1 \end{array} \right\rangle$$

Thus, we get from the relations that:

$$\begin{cases} \eta_3 = \alpha_1 \alpha_3 \alpha_1^{-1} \\ \eta_1 = \alpha_2 \alpha_1 \alpha_2^{-1} \\ \eta_2 = \alpha_3 \alpha_2 \alpha_3^{-1} \end{cases} \Rightarrow \begin{cases} \lambda_1 = \alpha_2^{-1} \alpha_3 \alpha_2 \alpha_3^{-1} \\ \lambda_2 = \alpha_3^{-1} \alpha_1 \alpha_3 \alpha_1^{-1} \\ \lambda_3 = \alpha_1^{-1} \alpha_2 \alpha_1 \alpha_2^{-1} \end{cases}.$$

And we can, finally, compute the Magnus expansion of each longitude, obtaining:

$$\begin{aligned} \Phi(\lambda_1) &= \Phi(\alpha_2^{-1} \alpha_3 \alpha_2 \alpha_3^{-1}) = (1 - t_2 + t_2^2)(1 + t_3)(1 + t_2)(1 - t_3 + t_3^2) \\ &= 1 - t_2 t_3 + t_3 t_2 \end{aligned}$$

Note that we have intentionally omitted the terms of degree greater than or equal to 3 since we will not be computing any Milnor's invariant with more than 3 components hence the terms with higher degrees will not be necessary.

Recall that any linking number of a unique component is by definition zero and, as we've seen before, the regular linking numbers are all also zero, thus we only need to compute all possible triple linking numbers. As before, we may compute the missing Magnus expansions for the other two longitudes to get:

$$\Phi(\lambda_2) = 1 - t_3 t_1 + t_1 t_3 \text{ and } \Phi(\lambda_3) = 1 - t_1 t_2 + t_2 t_1.$$

Then, we get:

$$\mu_B(12, 3) = -1 \Rightarrow \bar{\mu}_B(123) = -1$$

$$\mu_B(3\,1,2) = -1 \Rightarrow \bar{\mu}_B(3\,1\,2) = -1$$

$$\mu_B(2\,3,1) = -1 \Rightarrow \bar{\mu}_B(2\,3\,1) = -1$$

$$\mu_B(1\,3,2) = 1 \Rightarrow \bar{\mu}_B(1\,3\,2) = 1$$

$$\mu_B(2\,1,3) = 1 \Rightarrow \bar{\mu}_B(2\,1\,3) = 1$$

$$\mu_B(3\,2,1) = 1 \Rightarrow \bar{\mu}_B(3\,2\,1) = 1$$

Hence, the three components of B are tangled even though any two of them is not. This is the precise characteristic of the Borromean rings that we were pursuing to prove.

3 The First Quantum Invariants

We have already studied the linking number, which is a very important knot invariant, we have also seen a little bit about the knot group and how this also works out to be a knot invariant. Now we start to look at some extremely powerful tools that are highly used in the study of knots: the Jones polynomial and the Kauffman bracket. The Jones polynomial as we will see, appears not only on the study of classification of knots, which is one of the main goals on knot theory, but it also appears on physics, when we study for example the world line, a representation on how things move throughout space as time passes, we will see briefly how this kind of studies reveals a familiar object that can be topologically studied by means of the Jones polynomials.

3.1 The Kauffman Bracket

Let's remember that a Laurent polynomial is a polynomial written in the following way:

$$p(t) = \dots + a_{-2}t^{-2} + a_{-1}t^{-1} + a_0 + a_1t + a_2t^2 + \dots$$

thus, we are going to determine a way of relating knot diagrams to Laurent polynomials with integer coefficients, on the variable A , i.e., if $\mathcal{D}$ is set of all diagrams of non oriented links, then our objective here is to define functions $\mathcal{D} \rightarrow \mathbb{Z}[A^{-1}, A]$.

The first way of doing this, is by using the Kauffman bracket:

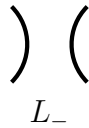
Definition 3.1.1. The *Kauffman bracket* is the only function in $\mathbb{Z}[A^{-1}, A]$ defined in $\mathcal{D}$ that satisfies the following properties:

$$\langle \bigcirc \rangle = 1$$

$$\langle D \sqcup \bigcirc \rangle = (-A^{-2} - A^2) \langle D \rangle$$

$$\langle L_0 \rangle = A \langle L_+ \rangle + A^{-1} \langle L_- \rangle$$

Where, in Definition 3.1.1, $\bigcirc$ represents the diagram of the unknot, and L_0 , L_+ and L_- are the crossings:



The last relation on Definition 3.1.1 is called the *Skein relation* for the Kauffman bracket. Also, the squared union symbol $\sqcup$ means that the union between the diagrams D and $\bigcirc$ is disjoint, on the second relation. From now on, unless we tell differently we are going to be using these introduced notations.

Observe now that given $D \in \mathcal{D}$, if we apply successively the Skein relations of the Kauffman bracket on D , we can write $\langle D \rangle$ as $\langle \bigcirc \sqcup \dots \sqcup \bigcirc \rangle$, thus $\langle D \rangle$ is a linear combination of brackets of disjoint unions of the trivial link $\bigcirc$, which can be easily calculated by using the first two relations on the definition of the Kauffman bracket.

If we take $\overline{D}$ to be the diagram of the link D but with all the crossings changed (meaning that if on D a crossing is of the type of L_0 , then on $\overline{D}$, the same crossing is now L_1). Also, if we denote $\langle D \rangle^*$ by changing to negative all positive exponents of $\langle D \rangle$ and to positive all negative exponents, i.e., $*$: $\mathbb{Z}[A^{-1}, A] \rightarrow \mathbb{Z}[A, A^{-1}]$ is given by $*(A) = A^{-1}$ and $*(A^{-1}) = A$, then we have that the following equality holds:

$$\langle \overline{D} \rangle = \langle D \rangle^*$$

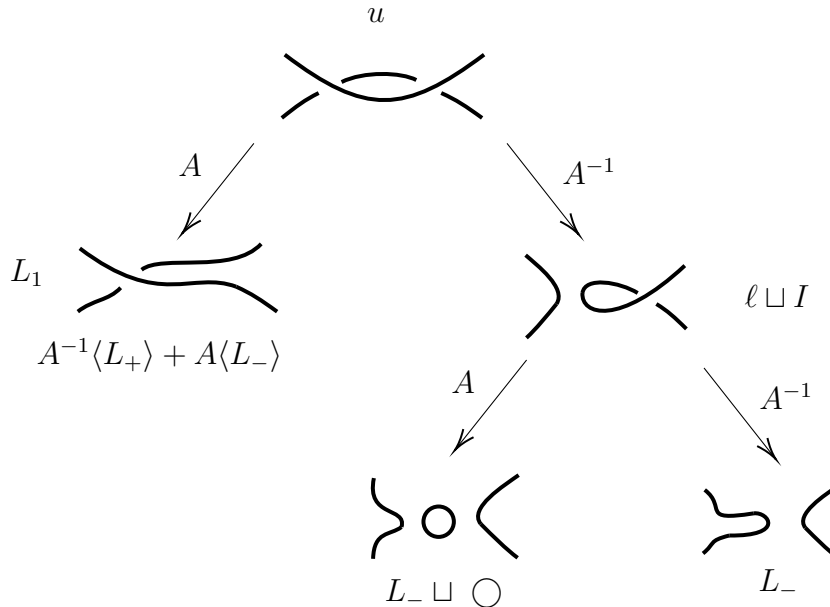
And by using this, we can have a fourth relation for the Kauffman bracket that comes directly from the third one already given:

$$\langle L_1 \rangle = A^{-1} \langle L_+ \rangle + A \langle L_- \rangle.$$

An important thing we would like to happen is that the Kauffman bracket is a link invariant, unfortunately, the following proposition only guarantee that it is invariant by two out of three Reidemeister moves.

Proposition 3.1.1. *The Kauffman bracket is invariant by Reidemeister moves R_2 and R_3 .*

Proof. We must check that $\langle u \rangle = \langle R_2 u \rangle$ and $\langle X \rangle = \langle R_3 X \rangle$. There is no way here to run away from using diagrams, so let's first do a picture of what is happening in each case diagrammatically at the same time we introduce the notations, then we rewrite everything in terms only of algebraic notations. First, let's work with the R_2 move. Notice that we can begin with any crossing on the u diagram, after applying all the relations we will arrive at the same result:

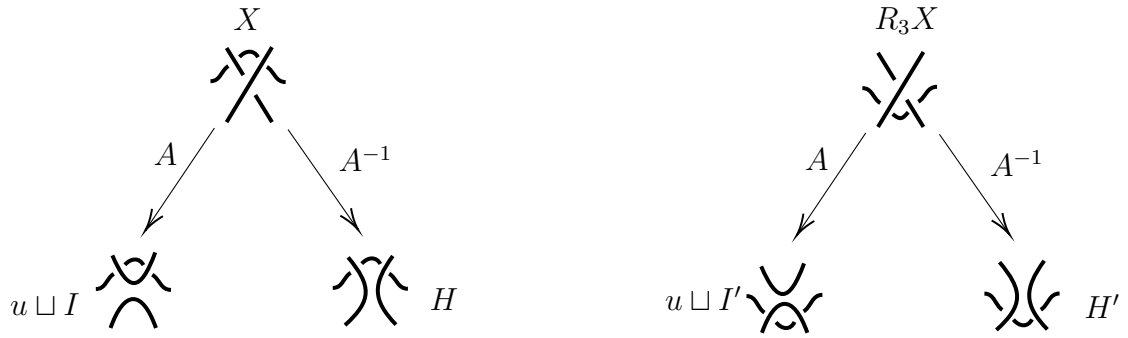


Therefore, the Kauffman bracket of u is given by:

$$\begin{aligned}
 \langle u \rangle &= A \langle L_+ \rangle + A^{-1} \langle \ell \sqcup I \rangle \\
 &= A \left(A^{-1} \langle L_+ \rangle + A \langle L_- \rangle \right) + A^{-1} \left(A \left(-A^{-2} - A^2 \right) \langle L_- \rangle + A^{-1} \langle L_- \rangle \right) \\
 &= \langle L_+ \rangle + A^2 \langle L_- \rangle + A^{-1} \left(-A^{-1} - A^3 \right) \langle L_- \rangle + A^{-2} \langle L_- \rangle \\
 &= \langle L_+ \rangle + \cancel{A^2 \langle L_- \rangle} - \cancel{A^{-2} \langle L_- \rangle} - \cancel{A^2 \langle L_- \rangle} + \cancel{A^{-2} \langle L_- \rangle} \\
 &= \langle L_+ \rangle
 \end{aligned}$$

Notice, that in this case since the diagrams are non oriented we have $R_2 u = L_+$, hence $\langle u \rangle = \langle L_+ \rangle = \langle R_2 u \rangle$.

Now, for the R_3 move observe the following diagrams:



Observe also that $u \sqcup I$ and $u \sqcup I'$ are the same diagram, just rotated π radians, and H and H' are the same diagram, also rotated π radians, therefore they have the same Kauffman bracket, hence $\langle X \rangle = \langle R_3 X \rangle$, and that completes the proof. $\square$

Although the theorem is true for Reidemeister moves R_2 and R_3 , it does not hold for Reidemeister move R_1 , in fact, a simple computation shows that $\langle \ell \rangle \neq \langle R_1 \ell \rangle$:

$$\begin{aligned}
 \langle \ell \rangle &= A^{-1} \langle R_1 \ell \rangle + A \langle R_1 \ell \sqcup \bigcirc \rangle \\
 &= A^{-1} \langle R_1 \ell \rangle + A \left((-A^{-2} - A^2) \langle R_1 \ell \rangle \right) \\
 &= A^{-1} \langle R_1 \ell \rangle - A^{-1} \langle R_1 \ell \rangle - A^3 \langle R_1 \ell \rangle \\
 &= -A^3 \langle R_1 \ell \rangle
 \end{aligned}$$

We here, are considering ℓ exactly like established in Theorem 2.2.1, but the term $-A^3$, found above, can change, by changing the type of the crossing in ℓ , in fact, if $\bar{\ell}$ denotes the diagram ℓ with the crossing changed, then $\langle \bar{\ell} \rangle = -A^{-3} \langle R_1 \bar{\ell} \rangle$. Although the Kauffman bracket is not a knot invariant, since it is not invariant through the R_1 move, it is fundamental in the definition of a, soon to be defined, invariant. But first we need to look at one more property of knots, that will allow us to define this new invariant.

3.2 The Writhe and the Jones Polynomial

While the linking number gives us a way to measure "how linked two links are", the writhe will give us a way to see how much is a link twisted around itself.

Definition 3.2.1. Let $\vec{D}$ be a oriented link, we define the *writhe* of $\vec{D}$ as the sum of all signs $\varepsilon(p)$ for all crossings p of the diagram $\vec{D}$, we denote the writhe of $\vec{D}$ as $w(D)$ or sometimes $w(\vec{D})$.

In the definition above, the sign of a crossing is defined in the same fashion we defined the sign on the linking number, since the diagram is oriented. It is important to notice here, that the writhe varies over *all* crossings of a diagram, differently from the linking number, where the crossings we were interested in were the crossings between different components.

Example. Consider the diagram of the trefoil given in Section 2.6, Figure 3.3, and call it $\vec{T}$. Now, observe that in this diagram we have three different crossings, a simple observation on the signs of this crossings gives us $w(T) = -1 - 1 - 1 = -3$.

Example. Consider a diagram of the unknot, with any orientation. Since it has no crossings, its writhe is zero.

Observe now that taking any orientation for the diagram u , then $w(u) = w(R_2u)$ and the same happens for X and R_3X , this follow in the same way we did for the linking number, since the definition of the sign for the linking number is the same as the one used to define the writhe. Now, on the other hand note that if ℓ is oriented from left to right (ℓ as in Theorem 2.2.1), let us denote this by $\vec{\ell}$ (and if ℓ is oriented from right to left we denote $\overleftarrow{\ell}$), then we have

$$w(\vec{\ell}) = w(R_1\vec{\ell}) + 1 \quad \text{and} \quad w(\overleftarrow{\ell}) = w(R_1\overleftarrow{\ell})$$

where $R_1\vec{\ell}$ is the diagram $R_1\ell$ with orientation compatible with that of $\vec{\ell}$ and similarly for $R_1\overleftarrow{\ell}$.

Thus, the writhe is also not an invariant for the R_1 move. But, fortunately, a specific combination of both the writhe and the Kauffman bracket, gives us a nice knot invariant, i.e., a property of each link that is invariant through all three Reidemeister moves, and this is the so called *Jones polynomial*, stated in the following theorem:

Theorem 3.2.1. Let $\vec{D}$ be a diagram for the oriented link L and D denotes the non oriented adjacent diagram. Then the function

$$V_L = (-A)^{-3w(\vec{D})} \langle D \rangle \in \mathbb{Z}[A^{-1}, A]$$

where $w(\vec{D})$ is the writhe of $\vec{D}$, is a link invariant, called the Jones polynomial of L .

Proof. Since the writhe and the Kauffman bracket are both invariants of R_2 and R_3 moves, then the Jones polynomial is clearly invariant of the same moves.

Now, for the R_1 move, taking $\langle \ell \rangle = -A^3 \langle R_1\ell \rangle$, the associated writhe for this diagram is $w(\vec{\ell}) = w(R_1\vec{\ell}) + 1$ (remember we have two possibilities of orientation, the second one follows in the same way), we have

$$\begin{aligned} V_L(\vec{\ell}) &= \left((-A)^{-3w(\vec{\ell})} \right) \langle \ell \rangle = \left[(-A)^{-3w(R_1\vec{\ell})-3} \right] (-A^3) \langle R_1\ell \rangle \\ &= (-A)^{-3(R_1\vec{\ell})} \langle R_1\ell \rangle = V_L(R_1\ell) \end{aligned}$$

Hence, the Jones polynomial is a link invariant. $\square$

Notice now that $V_{\bigcirc} = A^{-3w(\bigcirc)}\langle\bigcirc\rangle = A^{-3\cdot 0} \cdot 1 = 1$, for any orientation of $\bigcirc$. Furthermore, even though we are using orientation to define the writhe, if we change the orientation of all the components of a link, its writhe will not change.

However, let L be a link and K a sublink of L , let $L - K$ be the link L with every component of K removed, and denote by $r_K(L)$ the link L with all the orientations of K reversed, then

$$V_{r_K(L)} = A^{12lk(K, K-L)} V_L.$$

Also, another important thing to mention about the Jones polynomial is that it is more common to use it in the variable t , by doing the substitution $A^{-2} = t^{\frac{1}{2}}$, thus, $V_L \in \mathbb{Z}[t^{-\frac{1}{2}}, t^{\frac{1}{2}}]$.

Now, we introduce an important result, it tells us that the Jones polynomial is the only polynomial satisfying a certain relation, also called the skein relation.

Theorem 3.2.2. *The Jones polynomial is the only link invariant satisfying the following Skein relation:*

$$t^{-1}V_{L_1} - tV_{L_0} = (t^{\frac{1}{2}} - t^{-\frac{1}{2}}) V_{L_-}$$

with the normalization $V_{\bigcirc} = 1$, where $\bigcirc$ represents the unknot oriented clockwise. In particular, $V_L \in \mathbb{Z}[t^{\frac{1}{2}}, t^{-\frac{1}{2}}]$.

Proof. First, note that:

$$\begin{aligned} A\langle L_1 \rangle - A^{-1}\langle L_0 \rangle &= A^2\langle L_- \rangle + \langle L_+ \rangle - A^{-2}\langle L_- \rangle - \langle L_+ \rangle \\ &= (A^2 - A^{-2})\langle L_- \rangle \end{aligned}$$

Furthermore, $w(L_1) = w(L_-) + 1$ and $w(L_0) = w(L_-) - 1$. Thus, computing $A^4V_{L_0} - A^{-4}V_{L_1}$, making use of Theorem 3.2.1, we get:

$$A^4V_{L_0} - A^{-4}V_{L_1} = (A^2 - A^{-2}) \left[(-A^{3w(L_-)})\langle L_- \rangle \right] = (A^2 - A^{-2})V_{L_-}.$$

On the other hand, notice that any link can be transformed into the trivial link with k components if we allow change of crossings, thus it is enough to find $V_{k\bigcirc}$, for the trivial link with k components $k\bigcirc$. Note that:

$$V_{\bigcirc} = 1, \quad V_{\bigcirc \sqcup \bigcirc} = (-A^2 - A^{-2}), \dots, V_{k\bigcirc} = (-A^2 - A^{-2})^{k-1}$$

which in terms of t , we have $V_{k\bigcirc} = (t^{-\frac{1}{2}} - t^{\frac{1}{2}})^{k-1}$. □

It is also important to observe that, before the Jones polynomial was invented, another (really similar) invariant was known to satisfy a similar Skein relation, this invariant, was the Alexander polynomial. Which is a polynomial $\nabla_L \in \mathbb{Z}[t, t^{-1}]$ satisfying:

$$-\nabla_{L_1} - \nabla_{L_0} = (t - t^{-1})\nabla_{L_-}.$$

Although it looks like the Alexander polynomial satisfies the same skein relation as that of the Jones polynomial, actually, the skein relation related to the Alexander polynomial occurs under certain normalization, called the Conway normalization.

3.3 Properties of the Jones Polynomial

The Jones polynomial have some advantages, or more precisely, some characteristics that its predecessor (the Alexander polynomial) does not have. We are talking about the capability of identifying chirality of links. The Jones polynomial is capable, thanks to the way it is defined, to identify when a link is equivalent to its mirror image. And this property of being equivalent to its mirror image has a name:

Definition 3.3.1. Let L be a link and $\bar{L}$ its mirror image. We say L is *amphichiral* when L is equivalent to $\bar{L}$.

Let D be the diagram of a link L and $\bar{D}$ the diagram of $\bar{L}$. Remember, we stated, when talking about the Kaffman bracket, that $\langle \bar{D} \rangle = \langle D \rangle^*$, also, note that $w(L) = w(\bar{L})$, since reflecting L does not change its crossings. Thus:

$$V_D(t) = V_{\bar{D}}(t).$$

For example, take T the diagram for the trefoil knot, and H the diagram for the Hopf link, with orientation compatible with that of the trefoil after smoothing one crossing of T . Then we get

$$t^{-1}V_T(t) - tV_{\bigcirc}(t) = (t^{1/2} - t^{-1/2})V_H(t) \Rightarrow t^{-1}V_T(t) - t = (t^{1/2} - t^{-1/2})V_H(t) \quad (*)$$

Also, now looking at the diagram of H we see that:

$$\begin{aligned} t^{-1}V_H(t) - tV_{\bigcirc \sqcup \bigcirc}(t) &= (t^{1/2} - t^{-1/2})V_{\bigcirc}(t) \Rightarrow t^{-1}V_H(t) - t(-t^{1/2} - t^{-1/2}) = t^{1/2} - t^{-1/2} \\ &\Rightarrow t^{-1}V_H(t) + t^{3/2} + t^{1/2} = t^{1/2} - t^{-1/2} \\ &\Rightarrow V_H(t) = -t^{5/2} - t^{1/2} \end{aligned}$$

And now, replacing $V_H(t)$ on $(*)$, we obtain:

$$\begin{aligned} t^{-1}V_T(t) - t &= (t^{1/2} - t^{-1/2})(-t^{5/2} - t^{1/2}) \Rightarrow t^{-1}V_T(t) = -t^3 + t^2 + 1 \\ &\Rightarrow V_T(t) = -t^4 + t^3 + t. \end{aligned}$$

And since this polynomial is not an invariant through $t \mapsto t^{-1}$ (just note that the image of V_T , would be $V_{\bar{T}}(t) = -t^{-4} + t^{-3} + t^{-1} \neq V_T(t)$), then, T does not have the aphichiral property.

A similar computation, shows us that, contrary to the trefoil, the figure eight knot is aphichiral, since its Jones polynomial is given by $V_E(t) = t^{-2} - t^{-1} + 1 - t + t^2$.

3.4 Names of knots and the Jones polynomial table

In this section, we will see how we name knots, from the already known ones like the trefoil, and the figure eight, to knots we haven't seen yet. Also, we will establish a table giving the coefficients of the Jones polynomial for the first knots. The computation for all this Jones polynomials have been done by lots of mathematicians and shall not be redone here.

Let's first introduce the notation for the trefoil knot, which is 3_1 , this notation comes from the fact that the trefoil have 3 crossings, and is the first diagram with 3 crossings, in fact, the trefoil is the only diagram with 3 crossings, meaning that every other diagram with 3 crossings, must be equivalent to the trefoil thus it shall not be considered on the table. Also, note that the maximum sub-index of 3 will be 1. It is obvious that we do not have any knot with one or two crossings that is not the unknot, therefore our table shall start with the trefoil. For 4 crossings, the first diagram we can think of is the figure eight knot, and in fact, in the same fashion as the trefoil, any other diagram with four crossings must be a diagram for the figure eight knot, therefore, the family of diagrams with 4 crossings have only one element, and thus the name for the figure eight will be 4_1 .

For 5 crossings, we have 2 knots that are not equivalent although they have the same amount of crossings (which is 5), so every diagram with 5 crossings must be equivalent to 5_1 or to 5_2 , which are the two first (and only) distinct diagrams with 5 crossings.

Following this line of thought, we have 3 distinct diagrams with 6 crossings; 7 distinct diagrams with 7 crossings; 21 distinct diagrams with 8 crossings, and so on. Now, the following page contains a table with the coefficients of the Jones polynomial of each of the knots from 3_1 to 8_{21} . The bold number is the coefficient of the term t^0 on the Jones polynomial, and we are considering polynomials of the form $\dots + a_{-2}t^{-2} + a_{-1}t^{-1} + a_0t^0 + a_1t + a_2t^2 + \dots$

Knot											
3 ₁	-1	1	0	1	0						
4 ₁	1	-1	1	-1	1						
5 ₁	-1	1	-1	1	0	1	0	0			
5 ₂	-1	1	-1	2	-1	1	-1	0			
6 ₁	1	-1	1	-2	2	-1	1				
6 ₂	1	-2	2	-2	2	-1	1				
6 ₃	-1	2	-2	3	-2	2	-1				
7 ₁	-1	1	-1	1	-1	1	0	1	0	0	0
7 ₂	-1	1	-1	2	-2	2	2	-1	1	0	
7 ₃	0	0	1	-1	2	-2	3	-2	1	-1	
7 ₄	0	1	-2	3	-2	3	-2	1	-1		
7 ₅	-1	2	-3	3	-3	3	-1	1	0	0	
7 ₆	-1	2	-3	4	-3	3	-2	1			
7 ₇	-1	3	-3	4	-4	3	-2	1			
8 ₁	1	-1	1	-2	2	-2	2	-1	1		
8 ₂	1	-2	2	-3	3	-2	2	-1	1		
8 ₃	1	-1	2	-3	3	-3	2	-1	1		
8 ₄	1	-2	3	-3	3	-3	2	-1	1		
8 ₅	1	-1	3	-3	3	-4	3	-2	1		
8 ₆	1	-2	3	-4	4	-4	3	-1	1		
8 ₇	-1	2	-2	4	-4	4	-3	2	-2		
8 ₈	-1	2	-3	5	-4	4	-3	2	-1		
8 ₉	1	-2	3	-4	5	-4	3	-2	1		
8 ₁₀	-1	2	-3	5	-4	5	-4	2	-1		
8 ₁₁	1	-2	3	-5	5	-4	4	-2	1		
8 ₁₂	1	-2	4	-5	5	-5	4	-2	1		
8 ₁₃	-1	2	-3	5	-5	5	-4	3	-1		
8 ₁₄	1	-3	4	-5	6	-5	4	-2	1		
8 ₁₅	1	-3	4	-6	6	-5	5	-2	1	0	0
8 ₁₆	-1	3	-5	6	-6	6	-4	3	-1		
8 ₁₇	1	-3	5	-6	7	-6	5	-3	1		
8 ₁₈	1	-4	6	-7	9	-7	6	-4	1		
8 ₁₉	0	0	0	1	0	1	0	0	-1		
8 ₂₀	-1	1	-1	2	-1	2	-1				
8 ₂₁	1	-2	2	-3	3	-2	2	0			

Figure 3.1: Table of Jones Polynomials.¹

¹This table is based on a table from "An Introduction to Knot Theory" from W. B. Raymond Lickorish
- page 27.

3.5 Mutation

An issue we encounter when working with the Jones polynomial is that there exists infinite many pairs of distinct knots with the same Jones polynomial. And a process to find such pairs of knots is through a concept called *mutation*. We can describe the mutation as follows: consider L a link and take D^3 the 3-disk that intersects L in exactly 4 points in its boundary, i.e., $\partial D^3 \cap L = \{x_1, x_2, x_3, x_4\}$. Now choose a rotation axis for D^3 in such a way that a rotation of π rad keeps $\partial D^3 \cap L$ fixed.

Now, excising $D^3 \cap L$ from L and replacing it by $R(D^3 \cap L)$, where R is such rotation, then we obtain a mutation of the link L by the rotation R of D^3 . The following picture represents two knots that are related by a mutation, when that occurs, we say the knots are *mutant*. Figure 3.2 shows two mutant knots, the Kinoshita-Terasaka (11n42) and the Conway knot (11n34):

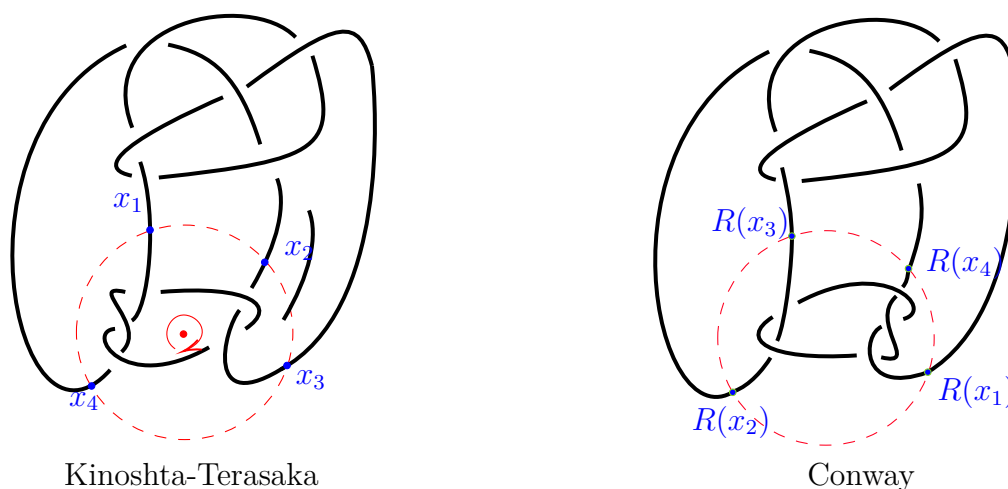


Figure 3.2: Two mutant knots

The next proposition, gives us the exactly problem we started the section, we will see that any pair of mutant knots have the same Jones polynomial, hence, it turns out that the Jones polynomial is not a extremely good tool to identify different knots.

Proposition 3.5.1. *Two knots obtained by means of a mutation have the same Jones polynomial.*

Proof. We can use the skein relation on non oriented diagrams to write down a formal sum of diagrams for $D^3 \cap L$ as follows:

$$\alpha L_- + \beta L_+$$

this expression does not change by any rotation of $D^3 \cap L$ that keeps fixed $\partial D^3 \cap L$, since in this decomposition, the crossings are smoothed out. Now, note that the writhe of each diagram in $D^3 \cap L$ is the same, since they have the same crossings. Also, observe that if L' is the link obtained from L by a mutation, then to compute the Kauffman bracket of each knot, we first smooth the crossings of $D^3 \cap L$ out, but the formal sums of the smoothed crossings are the same, and also since the links are the same on $L - D^3 \cap L$, the Kauffman bracket of each link is the same.

Thus, since $w(L) = w(L')$ and $\langle L \rangle = \langle L' \rangle$, the result easily follows. $\square$

Although equivalent knots have the same Jones polynomial, we see by this construction of mutant knots, that not all knots that have the same Jones polynomial, are equivalent. Therefore it is a good question to ask if all the knots that have trivial Jones polynomial is the unknot. Unfortunately, this is still an answered question, although it looks like this is true, there have been found no knots non equivalent to the unknot that have trivial Jones polynomial, nor a proof for this statement has been done.

However, some discoveries similar to this very question have been done, Eliahou, Kauffman and Thistlethwaite, have discovered a link which is not the trivial link, but it has the same Jones polynomial as the trivial link with 4 components:

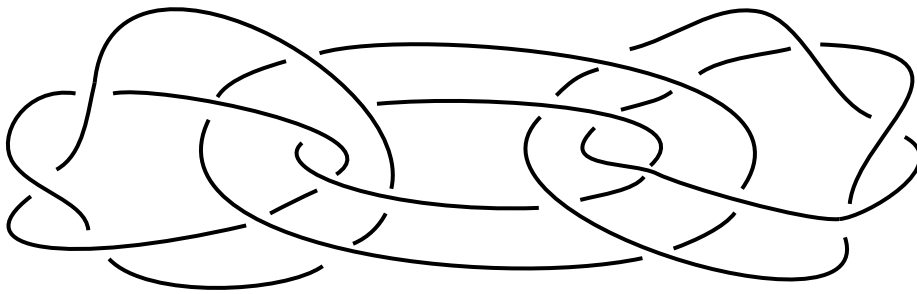


Figure 3.3: A non-trivial link, with trivial Jones polynomial

Another interesting result, is that there are also non mutant knots that have the same Jones polynomial, therefore knowing that mutant knots have the same Jones polynomial, does not mean that only those kinds of knots will have the same polynomial.

3.6 Alternating Links

We now turn our attention to a specific type of links, the *alternating links*. Alternating links can be defined in terms of its crossing types, more precisely, suppose you start at one crossing of the link, and travel through the link strand passing through all crossings, then in an alternating link, you should have, in sequence, crossings of the type over, under, over, under, and so on until you get back to the crossing you started. We can then give a definition for alternating links and diagrams.

Definition 3.6.1. A planar diagram is *alternating*, if the crossings we reach traveling through the components of the link are alternately overpasses and underpasses. A link is alternating if it has some alternating diagram.

The alternating links have some interesting properties that can be easily identified by looking at their diagrams, as we introduced in chapter 3, we will be able to identify primeness of a link and also connectedness of a link, just by studying their planar diagrams. We give now the formal definition for a link that is split, and an alternative definition for primeness of a link.

Definition 3.6.2. A link $L \subset S^3$ is called a *split link* if there exists a 2-sphere embedded in S^3 disjoint of L in such a way that each complementary ball of S^3 contains at least one component of the link L .

A diagram D of a link is *split* whenever a simple closed curve, embedded in $S^2 - D$ with each complementary disk containing at least one component of D , exists.

Definition 3.6.3. A non trivial link L is *prime* if for every 2-sphere $S^2 \subset S^3$ embedded in S^3 intersecting L in exactly two points, bounds on one side an arc satisfying the following: if we close this arc by gluing together the two points contained in the sphere, then the resulting knot is an unknot (we may simply say that S^2 bounds on one side an unknotted arc).

A non trivial diagram D is *prime* when each simple closed curve S^1 intersecting D in exactly two points bounds on one side a 2-disk intersecting D in a trivial arc.

Figure 3.4 shows exactly what the previous definitions are saying.

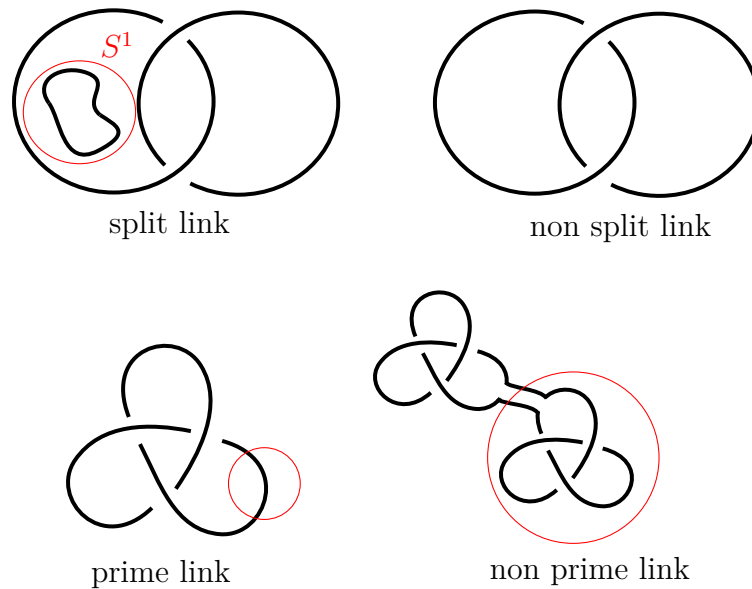


Figure 3.4: Prime, non prime, split, and non split links

It is important to observe that we can think of connectedness of links as the usual connectedness applied to the planar diagrams, but not considering the information about the crossings, if we were to consider such information, we would have to consider the great majority of knots as being disconnected, even the trefoil, for example, would be disconnected, which does not make sense, since if we think of the trefoil in $\mathbb{R}^3$, it is clearly connected. Thus it is common to refer to non split links as connected links.

Also, in the definition of primeness of a knot, the trivial arc bounded by the disk may be a (incomplete, thus an arc) diagram for the trivial knot with $n > 0$ crossings, it would of course imply that the knot is prime, however, if the arc has no crossings at all, i.e., the arc is an (incomplete) diagram for the trivial unknot (that without any crossings), then more than prime, we call this link *strongly prime*. As we will see, the condition of strongly primeness may be more important than it looks, since some results would not hold if this condition is even changed to only primeness.

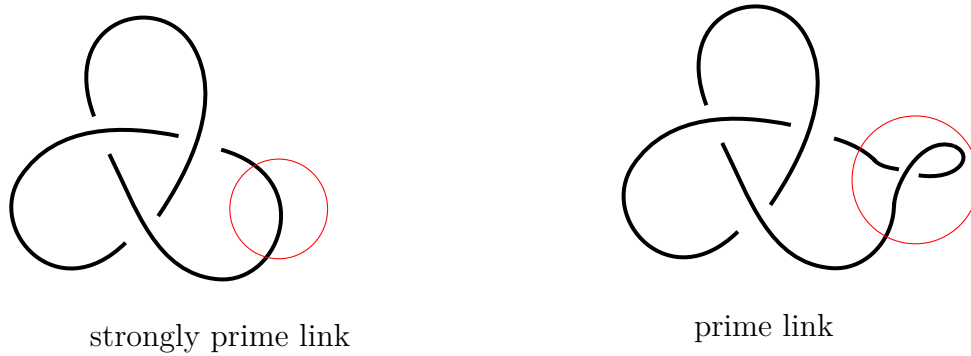


Figure 3.5: Prime and strongly prime links

A main problem on classification of knots is to determine whether a knot is split or not and if it is prime or not. The following result gives us one of the most useful tools for this section, which also facilitates a lot, the process of looking for non split and prime knots. The following proposition is proved in [15].

Proposition 3.6.1. *Let L be an alternating link with diagram D . Then L is split if, and only if D is split. And L is prime if, and only if D is prime.*

If we observe carefully, we are able to conclude that the first non trivial link that is non alternating, happens only with 8 crossings, meaning that every knot with less than 8 crossings that is non alternating is a diagram for the unknot. Hence, it is really easy identify primeness and splitness for links with less than 8 crossings.

If we think about it, there exists some kinds of crossings that can be undone. In fact, taking a trivial unknot and twisting it only once (giving us a link that looks like the symbol of infinity), it is obvious that we can twist it back, to where it was and going back to the trivial unknot. This kinds of crossings can appear in some link diagrams, and since we can "undo" them, reducing the number of crossings in the diagram by one, we call this kind of crossings a *reducible* crossings.

Definition 3.6.4. We say a crossing is *reducible* when there exists a closed simple curve intersecting the diagram D in exactly this crossing and no more points in the diagram.

Figure 3.6 can be thought of as each square representing a diagram of a link, which is omitted, but it can be anything. We are representing only the crossing that can be reduced:

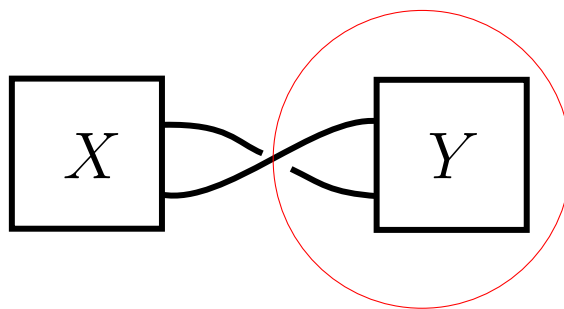


Figure 3.6: A reducible crossing

Another way of thinking on a reducible crossing is that if a crossing is reducible, then smoothing it out in a certain way, gives us a split link. That's what is shown in Figure 3.7

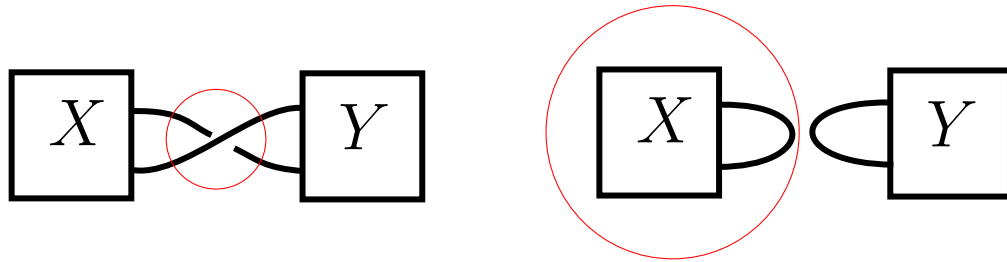


Figure 3.7: Another way of identify a reducible crossing

Definition 3.6.5. A *reduced diagram* is a diagram with no reducible crossings.

3.7 The Tait Conjectures

In the nineteenth century P. G. Tait, one of the founders of knot theory encountered some problems while trying to classify the knots. Since the study of knots at this time were lacking on mathematical rigor, and the Jones polynomial was yet to be discovered, the problems encountered by Tait remained as conjectures. The first two conjectures, which will be the focus on this section, were not proved until 1987 by Louis Kauffman, Kunio Murasugi and Morwen Thistlethwaite, and the proof, as we'll see, uses the Jones polynomial. Let's first state these two conjectures:

First conjecture: *"A reduced alternating diagram has the minimal number of crossings a diagram of the respective link can have".*

Second conjecture: *"The writhe of an alternating diagram is a topological invariant".*

We now introduce a important concept in the study of knots that is the *state* of a crossing on a diagram D . Remember that when we had an oriented diagram, we gave signs to each crossing depending on the type of the crossing. However, on a non oriented diagram, we can't even talk about signs of crossings, since the lack of an orientation stops us from identifying the signs. But since each crossing has no defined sign, we can attribute any sign to any crossing in any way we want. When we do that, we are defining a *state* for the crossing.

Definition 3.7.1. Suppose a diagram D has n crossings. If we label each crossing by the elements $\{1, 2, 3, \dots, n\}$, then the *state* of D is a function $s: \{1, 2, \dots, n\} \rightarrow \{-1, 1\}$. And to each sign we associate a smoothing of the diagram, if the crossing j have $s(j) = -1$ then we smooth it out as L_- , if $s(j) = 1$, then we smooth it out as L_+ .

Only for the sake of remembering the notation, we have the following picture for the smoothing:

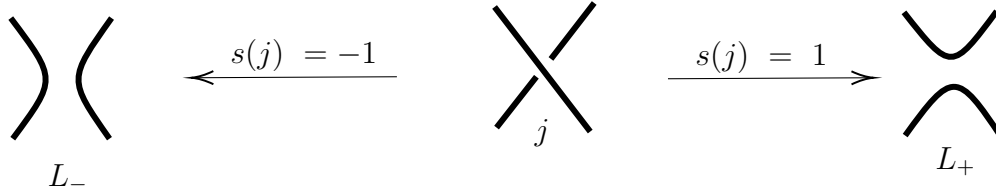


Figure 3.8: Smoothings of a state

Another way of see how to smooth a crossing is to put a label on it to indicate how the crossing should be smoothed out with respect to the given label, so we take again the same picture of Figure 3.8, but now we introduce the label on the crossing.

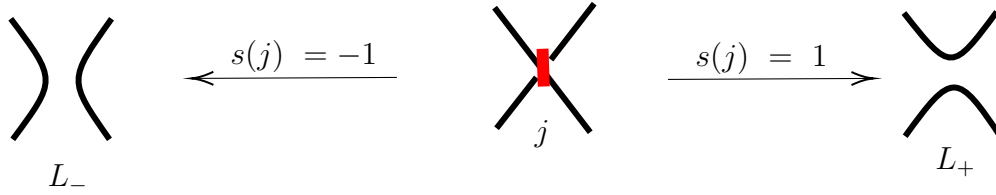
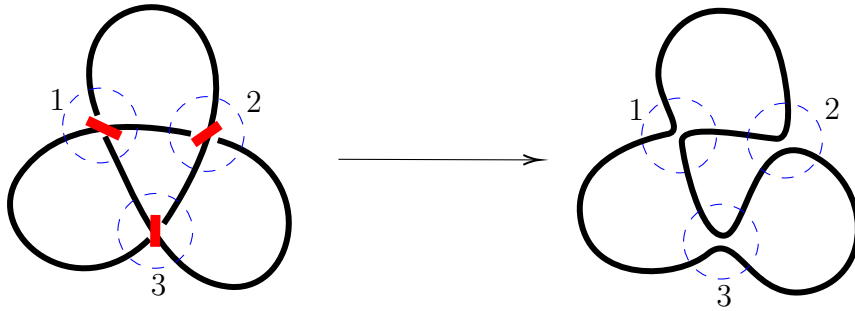


Figure 3.9: Smoothing of a state, with the notation of labels

The diagram obtained from D after the smoothing of all crossings is denoted as sD . This diagram is made of a finite disjoint union of circles (or simple closed curves), then the number of circles we obtain after a smoothing is denoted $|sD|$.

Example. Let's see how to work with the labels on a diagram by defining a state on a diagram of the trefoil T and smoothing the diagram with respect to the given state. Suppose we define the state on this diagram as $s(1) = 1$, $s(2) = -1$ and $s(3) = 1$, as the in Figure 3.10 below. Then, we have the following diagrams for the smoothing, and the labeled diagram:

Figure 3.10: State on the trefoil T

Now, consider the *breadth* of a Laurent polynomial V , denoted as $B(V)$ be the difference between the biggest exponent of V , denoted as $M(V)$ and the smallest exponent of V , denoted as $m(V)$. The main objective of this section will be to prove the following theorem:

Theorem 3.7.1. *Let D be a diagram of a connected (non split) link L with n crossings. If $V_L(t) \in \mathbb{Z}[t^{-1}, t]$ is the Jones polynomial of L , then:*

- (i) $B(V_L(t)) \leq n$;

- (ii) If D is reduced and alternating, then $B(V_L(t)) = n$;
- (iii) If D is non alternating and prime, then $B(V_L(t)) < n$.

But to prove this, we must first prove some auxiliary results. Maybe, the most interesting of them is the following lemma, which relates the Jones polynomial, with all the possible states of a diagram:

Lemma 3.7.1. *If D is a diagram of a link L with n crossings, then*

$$\langle D \rangle = \sum_s A^{\sum_{j=1}^n s(j)} (-A^2 - A^{-2})^{|sD|-1}.$$

Proof. Suppose $\langle D \rangle$ is the polynomial in $\mathbb{Z}[A, A^{-1}]$ satisfying:

$$\langle \bigcirc \rangle = 1$$

$$\langle D \sqcup \bigcirc \rangle = (-A^2 - A^{-2}) \langle D \rangle$$

$$\langle L_0 \rangle = A \langle L_- \rangle + A^{-1} \langle L_+ \rangle$$

then $\langle D \rangle$ is the Kauffman bracket of the diagram D . let's call $[D]$ the sum:

$$[D] = \sum_s A^{\sum_{j=1}^n s(j)} (-A^2 - A^{-2})^{|sD|-1}$$

and we will show that $[D]$ satisfies the same conditions of the Kauffman bracket, and since the Kauffman bracket is unique, we get $[D] = \langle D \rangle$.

- (i) For $\bigcirc$, since $s\bigcirc = \bigcirc$, then $|s\bigcirc| = 1$ and $\sum s(j) = 0$ since $\bigcirc$ have no crossings. Hence, we have that

$$[\bigcirc] = \sum_s A^{\sum s(j)} (-A^2 - A^{-2})^{|s\bigcirc|-1} = \sum_s A^0 (-A^2 - A^{-2})^{1-1} = 1.$$

- (ii) Note now that in doing $D \sqcup \bigcirc$, we haven't changed the value of the sum of the states, since $\bigcirc$ have no crossings, however we add one more $\bigcirc$ to the diagram sD , in other words, for $D \sqcup \bigcirc$, we have $\sum s(j)$ fixed and $|s(D \sqcup \bigcirc)| = |sD| + 1$. Thus,

$$\begin{aligned} [D \sqcup \bigcirc] &= \sum_s A^{\sum s(j)} (-A^2 - A^{-2})^{s(D \sqcup \bigcirc)-1} = \sum_s A^{\sum s(j)} (-A^2 - A^{-2})^{|sD|+1-1} \\ &= \sum_s A^{\sum s(j)} (-A^2 - A^{-2})^{|sD|-1} (-A^2 - A^{-2}) = (-A^2 - A^{-2}) [D]. \end{aligned}$$

- (iii) Let's now take D an arbitrary diagram and D_{L_-} another diagram that differs from D only on a crossing c . In this case note that D_{L_-} can be obtained from D , by defining on c the state $s(c) = -1$. In this case, we will have that, if u denotes the state on the diagram D_{L_-} and s , the state on D , we have that $\sum u(j) = (\sum s(j)) - 1$, moreover, the diagram sD is exactly the same as the diagram uD_{L_-} , since they are different in only one crossing, and in D this crossing has the state value of -1 . Thus, we have that

$$\begin{aligned} A[D_{L_-}] &= \sum_u A^{\sum u(j)+1} (-A^2 - A^{-2})^{|uD_{L_-}|-1} \\ &= \sum_{s, s(c)=-1} A^{\sum s(j)} (-A^2 - A^{-2})^{|sD|-1} \end{aligned} \tag{I}$$

In a similar fashion, for D_{L+} , we have that $\sum v(j) = \sum s(j) + 1$, with v a state for the diagram D_{L+} (a diagram that is different from D only on the crossing c , on which, D has the state value of $+1$, i.e., $s(c) = +1$). We have that

$$\begin{aligned} A^{-1}[D_{L+}] &= \sum_v A^{\sum v(j)-1} (-A^2 - A^{-2})^{|vD_{L+}|-1} \\ &= \sum_{s, s(c)=+1} A^{\sum s(j)} (-A^2 - A^{-2})^{|sD|-1} \end{aligned} \quad (\text{II})$$

Thus, note that adding (I) and (II), we have that the only difference between the terms, is that on the crossing c the first have the stated value of -1 while the second have the state value of $+1$, meaning that when adding the two terms together, we are considering all the possibilities of state for s and, hence it equals $[D]$, consequently,

$$[L_0] = A[L_-] + A^{-1}[L_+].$$

From (i), (ii) and (iii), follows that

$$\begin{aligned} [\bigcirc] &= 1 \\ [D \sqcup \bigcirc] &= (-A^2 - A^{-2})[D] \\ [L_0] &= A[L_-] + A^{-1}[L_+] \end{aligned}$$

thus, $[D]$ is the Kauffman bracket for the diagram D , i.e., $[D] = \langle D \rangle$, and the result follows. $\square$

Consider now the state s such that for every crossing j , $s(j) = 1$. This state is denoted by s_+ and we can call it the *positive state*, in the same way we can define the *negative state* as the state s_- such that $s_-(j) = -1$ for every crossing j on a diagram. Now observe the following facts about any state s of a diagram D :

- $\sum_{j=1}^n s(j) \equiv n \pmod{2}$. In fact, suppose s is positive in the set $\{1, \dots, k\}$ which means it is negative in $\{k+1, \dots, n\}$, then:

$$\sum_{j=1}^n s(j) = k - (n - k) = 2k - n \equiv n \pmod{2}$$

- $\left| \sum_{j=1}^n s(j) \right| = n \Leftrightarrow s = s_+ \text{ or } s = s_-$. In fact, if $\left| \sum_{j=1}^n s(j) \right| = n \Rightarrow s(j) = 1$, for all $n \in \{1, \dots, n\}$, or $\sum_{j=1}^n s(j) = -n \Rightarrow s(j) = -1$, for all $n \in \{1, \dots, n\}$, hence $s = s_+$ or $s = s_-$.

On the other hand, if $s = s_-$, then $\left| \sum_{j=1}^n s(j) \right| = \left| \sum_{j=1}^n (-1) \right| = |-n| = n$, since $n \in \mathbb{N}$. The same follow similarly for $s = s_+$.

- If we change the sign of only one crossing of D , for example a positive to a negative, then supposing the same amount of positive and negative signs as in the first item, then now we have $k-1$ positive crossings, and $n-k+1$ negative crossings, which gives us: $\sum_{j=1}^n s'(j) = (k-1) - (n-k+1) = 2k - n - 2 = \sum_{j=1}^n s(j) - 2$, where s' is the same as s except in the crossing we changed signs. A similar thought gives us that changing a negative sign to a positive sign will get us $\sum_{j=1}^n s'(j) = \sum_{j=1}^n s(j) + 2$.

Definition 3.7.2. Let D be the diagram of a knot. If $|s_+D| > |sD|$ for every state s such that $\sum_{j=1}^n s(j) = n - 2$, then D is called a *plus-adequate* diagram.

In other words, what this definition says is that, if we take all the states that differ from s_+ only by one crossing, then the amount of disjoint trivial knots after the smoothing of D through all these states will be less than the amount of disjoint trivial knots we get when smoothing the diagram through s_+ . Hence, when a diagram is plus-adequate we must have that for every crossing, the positive smoothing gives us parts of two distinct trivial knots.

Definition 3.7.3. Let D be the diagram of a knot. If $|s_-D| > |sD|$ for every state s such that $\sum_{j=1}^n s(j) = 2 - n$, then D is called a *minus-adequate* diagram.

In the same way as the plus-adequate diagram, if we take all the states that differ from s_- only by one crossing, then the amount of disjoint trivial knots after the smoothing of D through all these states will be less than the amount of disjoint trivial knots we get when smoothing the diagram through s_- . Hence, when a diagram is minus-adequate we must have that for every crossing, the negative smoothing gives us parts of two distinct trivial knots.

Definition 3.7.4. We call a diagram D *adequate*, when it is, at the same time, plus and minus-adequate.

Lemma 3.7.2. *A reduced diagram is adequate.*

Proof. Let's first note that proving this lemma is the same as proving that for any state s that is different from s_+ (or s_-) in exactly one crossing, when we smooth the diagram through s_+ then each strand of each L_0 crossing is contained in different components of s_+D .

Now, let's color the complement of the diagram D in a chessboard fashion. Then the components of s_+D are the boundary of monochromatic regions painted of one of the two colors used to paint the diagram complement, since the diagram is alternating. Similarly, the components of s_-D are boundaries for regions of the second color (i.e., if the components of s_+D bound regions of the first color, then the components of s_-D bound regions of the second color).

Taking any state s that is different of s_+ in only one crossing, we must have that in this crossing, the smoothing keeps the strands in the same component, meaning that in this crossing, there are two regions of the same color coming close to each other. Just note that since the diagram is irreducible, we must have that no colored region gets close to itself, if it does, the crossing where that occurs is reducible. Hence the number of components of sD is strictly smaller than the number of components of s_+D . And the same occurs in a similar way with s_- . Thus, the diagram is adequate. $\square$

Lemma 3.7.3. *Let D be a daiagram with n crossings of a link L . Then*

$$\begin{cases} M(\langle D \rangle) \leq n + 2|s_+D| - 2, & \text{with equality when } D \text{ is plus-adequate;} \\ m(\langle D \rangle) \geq -n - 2|s_-D| + 2, & \text{with equality when } D \text{ is minus-adequate.} \end{cases}$$

Hence, when the diagram is adequate we have:

$$B(\langle D \rangle) = 2n + 2|s_+D| + 2|s_-D| - 4.$$

Proof. Choose a sequence $s_0 = s_+, s_1, s_2, \dots, s_k = s$ where s_j changes the sign of the crossing i_j to -1 and $\{i_1, i_2, \dots, i_k\} \subset \{1, \dots, n\}$. Thus, $\sum_{j=1}^n s_r(j) = n - 2r$. Note that the diagrams $s_r D$ and $s_{r+1} D$ are the same except on the crossing i_r . Thus, $|s_r D| = |s_{r+1} D| \pm 1$. Then, we must have

$$M(\langle D|_{s_r} \rangle) - M(\langle D|_{s_{r+1}} \rangle) \in \{0, 4\},$$

i.e., $M(\langle D|_{s_r} \rangle) - M(\langle D|_{s_{r+1}} \rangle)$ and in particular, $M(\langle D|_{s_1} \rangle) \geq M(\langle D|_{s_{r+1}} \rangle)$. Note also that since $M(\langle D|_{s_+} \rangle) = n + 2|s_+ D| - 2$, then $M(\langle D \rangle) \leq n + 2|s_+ D| - 2$, which is exactly what we wanted.

Now, if D is plus-adequate, then for every state s , we have that $|s_1 D| = |s_+ D| \pm 1$, but since by hypothesis we have that $|s_+ D| > |s D|$, then $|s_1 D| = |s_+ D| - 1$. Hence, it follows that for every $s \neq s_+$, $M(\langle D|_s \rangle) < M(\langle D|_{s_+} \rangle)$ and by $M(\langle D|_{s_r} \rangle) - M(\langle D|_{s_{r+1}} \rangle)$, the sequence $s_0, s_1, \dots, s$ is decreasing and therefore, $M(\langle D \rangle) = M(\langle D|_{s_+} \rangle) = n + 2|s_+ D| - 2$.

For s_- the proof is completely analogous, but in this case, note that the sequence is $s_0 = s, s_1, \dots, s_k = s_-$ and this sequence is increasing. $\square$

Lemma 3.7.4. *if D is a connected diagram with n crossings, then $|s_+ D| + |s_- D| \leq n + 2$.*

Proof. For $n = 0$ the result follows at once, since there are no crossings, and therefore, $|s_+ \bigcirc| = |s_- \bigcirc| = 1$. Consider now a diagram with $n + 1$ crossings. Since D is connected by assumption, then selecting one crossing and smoothing it out in a positive or a negative way, we get a new connected diagram D' . Hence, supposing that this is obtained by the smoothing by s_- , then $|s_- D| = |s_- D'|$ (since $s_- D = s_- D'$) and also $|s_+ D| = |s_+ D'| \pm 1$, thus we have that

$$|s_+ D| + |s_- D| = |s_+ D'| + |s_- D'| \pm 1 \stackrel{(I.H.)}{\leq} n + 2 \pm 1 \leq n + 3.$$

$\square$

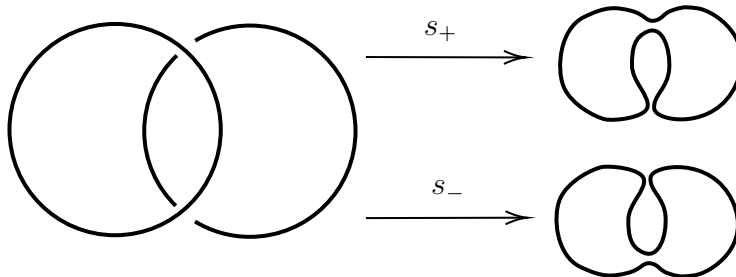
Lemma 3.7.5. *Let D be a connected diagram with n crossings*

(i) *If D is alternating, then $|s_+ D| + |s_- D| = n + 2$;*

(ii) *If D is non-alternating and strongly prime, then $|s_+ D| + |s_- D| < n + 2$.*

Proof. If D is alternating, by the proofs of the previous lemmas, we must have that if we color the complement of D in a chessboard fashion, then we get $|s_+ D|$ and $|s_- D|$ with the same number as the number of mono color components bounded by $s_+ D$ and $s_- D$ respectively. Thus $|s_+ D| + |s_- D|$ is the total number of these colored regions, and using the Euler characteristic, we get that this number is exactly $n + 2$.

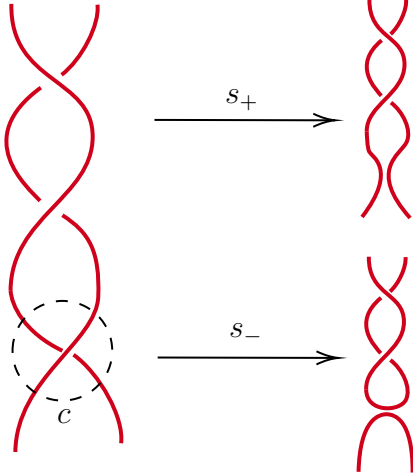
Now, suppose D is strongly prime and non-alternating. For $n = 2$, the result is immediate, since D is non alternating, the two crossings must have the same type, i.e., they are both over or both undercrossings, and since D is strongly prime, D must be the link with two components with two crossings being of the same type (this means that there exists some projection plane in which the link is disconnected), therefore:



$$|s_+D| + |s_-D| = 2 < 2 + 2 = 4.$$

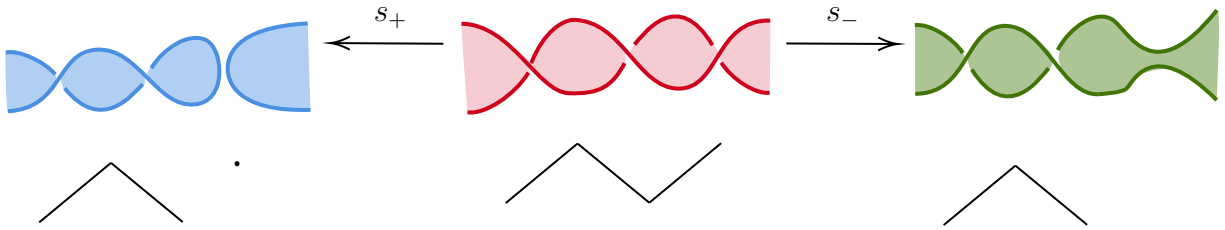
Notice the importance of the assumption of the link being strongly prime.

Now, suppose D has $n+1 \geq 3$ crossings. Then, there must be at least two consecutive crossings of the same type (again, both over or both overcrossings). Now, consider c a third crossing different from those two. Since D is strongly prime, both s_+ and s_- when applied to c gives us connected diagrams.



Again, considering the chessboard coloring of the complement of the diagram, we can associate a graph Γ to D , in which each vertex is a black region and the edges connect adjacent regions. Since D is strongly prime, then there is no loose point of Γ , i.e., a vertex that is separated from the rest of Γ . Thus, by removing this point, we keep the graph connected.

The two ways of smoothing this crossing c are represented by the following transformations on the graph:



If one way of smoothing the crossing out produces a vertex, the other produces a graph with no cut vertices. Thus, using the induction hypothesis we have that $|s_+D'| + |s_-D'| < 2 + n$, where D' is the strongly prime diagram formed by the previously smoothing. This implies that $|s_+D| + |s_-D| < n + 3$. $\square$

To conclude the proof of the first Tait's conjecture we first need to remember that $V_L(A) = (-A)^{-3w(\vec{D})}\langle D \rangle$, thus we have:

$$M_A(V_L) = M_A(\langle D \rangle) - 3w(\vec{D})$$

$$m_A(V_L) = m_A(\langle D \rangle) - 3w(\vec{D})$$

here we have put the variable A as a subindex since the initial theorem is stated in terms of the variable t , while the proofs we have given are in terms of the variable A , hence, we need to work with this change of variables. Recall also that $A^{-2} = t^{\frac{1}{2}}$. Hence, from the above, we get:

$$\begin{aligned} B_A(V_L) &= M_A(V_L) - m_A(V_L) = M_A(\langle D \rangle) - 3w(\vec{D}) - (m_A(\langle D \rangle) - 3w(\vec{D})) \\ &= M_A(\langle D \rangle) - m_A(\langle D \rangle) = B_A(\langle D \rangle). \end{aligned}$$

Now, from the change of variables we get:

$$\begin{aligned} -\frac{1}{4}M_A(V_L) &= m_t(V_L) \\ -\frac{1}{4}m_A(V_L) &= M_t(V_L) \end{aligned}$$

thus:

$$\begin{aligned} M_t(V_L) - m_t(V_L) &= -\frac{1}{4}m_A(V_L) + \frac{1}{4}M_A(V_L) = -\frac{1}{4}(-B_A(V_L)) \\ \Rightarrow B_t(V_L) &= \frac{B_A(V_L)}{4}. \end{aligned}$$

Putting together lemmas 3.7.2, 3.7.3 and 3.7.5(i) gives us that:

$$\begin{aligned} B_t(V_L) &= \frac{B_A(V_L)}{4} = \frac{B_A(\langle D \rangle)}{4} \\ &= \frac{1}{4} [2n + 2(|s_+D| + |s_-D|) - 4] \\ &= \frac{1}{4} (2n + 2n + 4 - 4) \\ &= n \end{aligned}$$

The other two items on the theorem follows in an analogous way, just using the corresponding equalities on each lemma. Thus we get that if a link L has a reduced alternating diagram D , then $B(V_L) = n$, thus taking any other diagram E for L , it can be alternating with m crossings, in which case following the theorem we get $B(\langle D \rangle) \leq n$, but since $B(\langle D \rangle) = n$, then $n \leq m$, and if the diagram E is non alternating, then $B(\langle D \rangle) < m$, and hence $n < m$. Finally, every alternating connected diagram of L have the minimal number of crossings a diagram of L can have.

3.8 The writhe of alternating links

The main objective of this section is to give an answer to the second Tait's conjecture, which, as we've seen, states that:

"The writhe of an alternating link is a topological invariant"

Before going to the main theorem which will allow us to deduce a positive answer to the conjecture, let's first establish the following lemma:

Lemma 3.8.1. *Let D be a diagram of an alternating link L . If D^r denotes a diagram formed by r copies of the diagram D , all parallel to each other and contained in the same plane, then if D is plus adequate, D^r is also plus-adequate.*

Proof. Note first that $(s_+D)^r = s_+(D^r)$ this can be easily seen by observing what happens to a crossing in the diagram D^r (which is actually r^2 crossings). Since each component of s_+D^r is parallel to a component of s_+D , then from the fact that D is alternating and plus-adequate we must have that none of the components of s_+D abuts itself around a former crossing, this would imply that after smoothing that crossing in a negative way, would give us an additional component, which contradicts the fact that D is plus-adequate. Hence, none of the components of $(s_+D)^r = s_+(D^r)$ abuts itself near a former crossing, therefore $s_+(D^r)$ is plus-adequate. $\square$

Theorem 3.8.1. Let $\vec{D}$ and $\vec{E}$ be two oriented diagrams of the same link L , having n_D and n_E crossings respectively. If D is plus-adequate, then

$$n_D - w(D) \leq n_E - w(E).$$

Proof. Let L_j be the various components of L , then D_j and E_j are the various diagrams of the corresponding components of L on the diagrams D and E respectively. Now, let's choose μ_j and ν_j positive integers, satisfying:

$$w(D_j) + \mu_j = w(E_j) + \nu_j.$$

Replacing each D_j by $\tilde{D}_j$, where $\tilde{D}_j$ is the component D_j with positive kinks, and also replacing E_j by $\tilde{E}_j$, where $\tilde{E}_j$ is constructed in the same fashion. Then the diagram $\tilde{D}$ formed by the $\tilde{D}_j$ components, remains plus-adequate, since the crossings on each positive kink, when smoothed out by s_+ , have strands in different components (which means the diagram is plus-adequate). Also, notice that, by construction we must have $w(\tilde{D}) = w(\tilde{E})$.

On the other hand, $\tilde{D}^r$ and $\tilde{E}^r$ are associated to the link L with r parallel copies similar to the construction of the D^r diagrams, by analogy, we can simply denote this as L^r . Then since $\tilde{E}^r$ and $\tilde{D}^r$ are diagrams for the same knot, they have the same Jones polynomial, but their writhe are also the same, thus $\langle \tilde{E}^r \rangle = \langle \tilde{D}^r \rangle$.

Now, note that, for each crossing on the diagram D , if $r = 2$ then, on the same crossing, D^2 have 4 new crossings, for $r = 3$, we have added 9 new crossings, and by induction it is easy to see that, for $r = n$ we have added n^2 new crossings. Also, note that for each r , after the smoothing by s_+ we obtain r new components. Hence, we must have the following:

$$\begin{aligned} M(\langle \tilde{E}^r \rangle) &\leq (n_E + \sum \nu_j) r^2 + 2(|s_+ E| + \sum \nu_j) r - 2 \\ M(\langle \tilde{D}^r \rangle) &= (n_D + \sum \mu_j) r^2 + 2(|s_+ D| + \sum \mu_j) r - 2 \end{aligned}$$

where the second equality comes from the fact that D is plus-adequate.

Therefore we must have

$$(n_D + \sum \mu_j) r^2 + 2(|s_+ D| + \sum \mu_j) r \leq (n_E + \sum \nu_j) r^2 + 2(|s_+ E| + \sum \nu_j) r$$

and since this holds for every r , we get $n_D + \sum \mu_j \leq n_E + \sum \nu_j$, and noticing that the writhe of D is the sum of the writhes of every D_j plus the linking number of L (the same for E), then by adding the linking number of L on both sides we get the final result:

$$n_D - w(D) \leq n_E - w(E).$$

□

Finally, to furnish an answer for the second Tait's conjecture, let's denote by $\#_D(+)$ the number of positive crossings of a diagram D for a link L , and by $\#_D(-)$ the number of negative crossings. Thus,

$$\begin{aligned} n_D &= \#_D(+) + \#_D(-) \\ w(\vec{D}) &= \#_D(+) - \#_D(-) \end{aligned}$$

hence, $n_D - w(\vec{D}) = 2\#_D(-)$ from what we can conclude, together with the previous result that:

- $\#_D(-) \leq \#_E(-)$ when D is plus adequate;
- $\#_D(+) \leq \#_E(+)$ when D is minus adequate.

and of course, if D is adequate, both items above hold at the same time, hence n_D must be minimal. Hence, two diagrams of the same link have the same writhe when they are both adequate.

4 The Braid Groups

The theory of braids was introduced Emil Artin in 1925. As we will see we have two distinct (however equivalent) ways of defining the braid groups, the first with a more algebraic perspective, and the second one a more intuitive and geometric point of view.

First, remember that two homeomorphisms are said to be isotopic, when there exists a homotopy h_t between them, with for every $t \in [0, 1]$, h_t is also a homeomorphism. Then we can define the *isotopy class* of a homeomorphism h as being the set of all homeomorphisms that are isotopic to h .

Consider $D = \{z \in \mathbb{R}^2; |z| \leq 1\}$ the closed unit disk in $\mathbb{R}^2$. Also, consider the set X consisting of $n \geq 1$ points in $\text{int}(D) \cap \mathbb{R} \times \{0\}$, meaning that the points x_i have coordinates $(x_i^1, 0)$ for every $i \in \{1, \dots, n\}$. Then, by calling D_n the set $D - X$, suppose there is a homeomorphism of D_n onto D_n , that is the identity when restricted to ∂D , when this happens, we say the homeomorphism preserves *pointwise* the set ∂D . Then the group $B_n = \pi_0(\text{Homeo}(D_n))$ consisting of all isotopy classes of those homeomorphisms of D_n that preserves pointwise the boundary ∂D_n , is called the *n-th Braid Group*. Now, taking a homeomorphism representative of a isotopy class of B_n , it extends uniquely to a homeomorphism $D \rightarrow D$ that permutes the points $x_1, x_2, \dots, x_n$. The class of such homeomorphism is called a *mapping class*¹. Also, this defines in a clear way, a homomorphism of groups between B_n and the symmetric group S_n .

Now, for $i = 1, 2, \dots, n$, the interval $[x_i^1, x_{i+1}^1] \subset (-1, 1) \subset \mathbb{R}$ is an embedded arc in D with its endpoints in the points x_i and x_{i+1} . Consider the homeomorphism that permutes counterclockwise the points x_i and x_{i+1} by a rotation of π radians and is the identity outside a disk with diameter $x_i^1 x_{i+1}^1$, such homeomorphism is clearly uniquely determined up to isotopy, and is called a *Dehn half twist* $D_n \rightarrow D_n$ that we denote by σ_i . Then the group B_n have generators $\sigma_1, \sigma_2, \dots, \sigma_{n-1}$ and satisfies the following relations:

$$\sigma_i \sigma_j = \sigma_j \sigma_i \text{ for } |i - j| \geq 2$$

$$\sigma_i \sigma_{i+1} \sigma_i = \sigma_{i+1} \sigma_i \sigma_{i+1} \text{ for } i = 1, 2, \dots, n - 2.$$

4.1 Some group notation

Now, let's remember that in a symmetric group a permutation that changes the position of only two elements is called a *transposition*. Suppose we have in S_4 the element that permutes 3 and 4 by changing their places, then, in S_4 the notation for this permutation will be

$$\begin{pmatrix} 1 & 2 & 3 & 4 \\ 1 & 2 & 4 & 3 \end{pmatrix}$$

¹The mapping class is a topological invariant

But, we can denote the same element by $(3\ 4)$, it is a transposition, and we can by this simple notation identify everything that the permutation is doing.

Another important thing to remember is that the set of transpositions

$$\{(1\ 2), (2\ 3), (3\ 4), \dots, (n-1\ n)\}$$

generates the symmetric group S_n .

Thus, we can take every element σ_i of B_n and relate it to $(i\ i+1)$. And this gives us the introduced homomorphism between these groups.

4.2 A geometric interpretation

Another interpretation of B_n is given in terms of *braids*, hence the name. First we define a *braid on n strings* as a n -component one dimensional manifold E contained in $D \times [0, 1]$ such that E intersects $D \times \{0, 1\}$ orthogonally on the set $(X \times 0) \cup (X \times 1)$ and the projection of each component of E maps homeomorphically onto $[0, 1]$. Geometrically, we see this notion of each component being taken onto $[0, 1]$ homeomorphically, as not allowing the braid to make a loop, like the loop on the Reidemeister move R_1 .

The braids are considered up to isotopy, therefore we can work, for example, with triangular moves to identify visually equivalent braids. Also, the braids are constant at the endpoints. The group operation is defined by gluing one braid under the other and compressing the result in $[0, 1] \times D$, as in the next figure

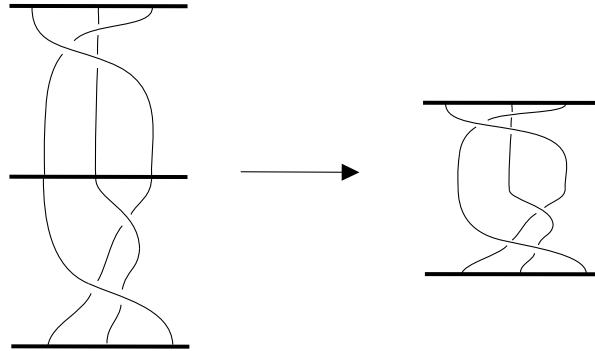


Figure 4.1: Operation on the braid group

4.3 Equivalence of definitions

Take a homeomorphism $h: D \rightarrow D$ that fixes pointwise the boundary of the disk. This homeomorphism is isotopic to the identity by a homotopy $h_t: D \times D \rightarrow D$, that can be represented by the family of homeomorphisms $\{h_t: D \rightarrow D\}_{t \in [0,1]}$, and such that $h_0 = Id_D$ and $h_1 = h$. If $h(X) = X$, then the union

$$\bigcup_{t \in [0,1]} (h_t(X) \times s) \subset D \times [0, 1]$$

is a braid.

We can imagine t as being the time, and the points x_i are particles moving around the disk D , and after $1s$ the particles are back to their initial points, but possibly permuted. Then as time passes, the graph of space versus time (which is a three dimensional graph)

represents a braid, showing all the paths along which each particle has moved while the time was passing. Each portion of this graph, i.e., if we fix an instant t_0 we can identify exactly where the particles were at that moment. Also note that since the movements are made by homeomorphisms, the particles do not just jump around, they all can be found somewhere around the plane, for every given instant t .

The isotopy class of each braid depends only on the element of B_n represented by h . This establishes an isomorphism between the B_n and the group of braids on n strings.

The generator of B_n correspond to the i -th elementary braid represented by a planar diagram consisting of n linear intervals which are disjoint except at one intersection point where the i -th interval goes over the $(i - 1)$ -th interval.

4.4 Alexander's theorem

We can associate to each braid, a link or a knot, which for generality we will only use the word link, but it must be clear that knots can also be associate to braids, and this will be clear in a second. To do that we need to define the closure of a braid.

Suppose a braid $x \in B_n$, as we have seen, the braid has an orientation, following the path of the homeomorphism that defines it, usually we take that orientation as being from up to down on the planar diagram of the braid. Now suppose from each point x_i on the domain D we extend an arc going outside the diagram of the braid until the i -th position on D on the image, note that the i -th position may not be the image of x_i . Then by doing this for each i , we get the planar diagram of a link that inherits the orientation from the braid, as in the following Figure.

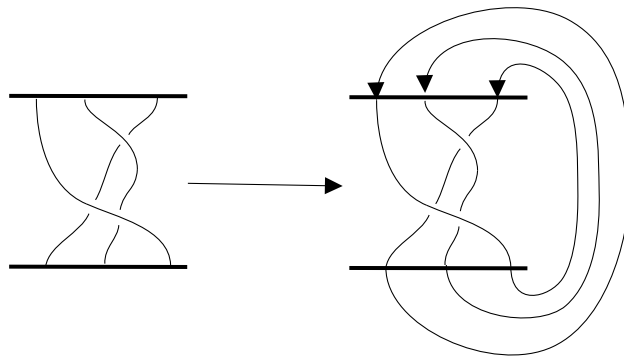


Figure 4.2: The closure of a braid

Note that the result is a link with two components, in fact, if we start at x_1 (as being the first string from left to right on the braid) and travel along the string we arrive at the beginning of the third string x_3 on the braid, and if we continue, we get back to x_1 , thus, it forms a component of the link. Also, if we start at the middle string, x_2 , we arrive again at x_2 after completing the string on the closure, hence it makes a new component of the link. Therefore, by completing our argument that there are braids that are associated to knots when closed, just imagine a braid composed by x_1 and x_3 on the previous picture, since it is a component of a link, it is a knot.

This happens because we are permuting the points x_1 and x_3 in some part of the braid, hence the "outside" strings (those formed on the closure) are not being connected

to themselves, instead, they are being connected to one another forming a new "bigger" string, a knot.

Now, the question that naturally arises is: can we always associate a link to a braid? And the answer is given as a theorem, the so called Alexander's theorem. James Waddell Alexander II, who was the first to note that the answer to our question is positive, however, the correspondence between links and braids is not a bijection, i.e., there are links that can be represented by to different braids.

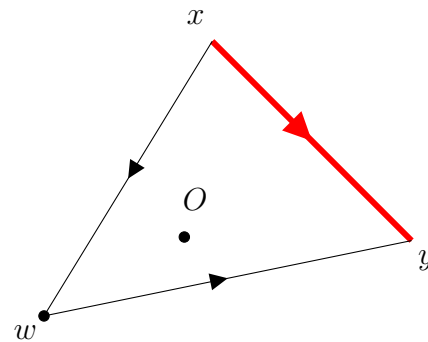
Theorem 4.4.1. *Any oriented link in $\mathbb{R}^3$ is isotopic to a closed braid.*

Proof. Consider D an oriented link diagram. Suppose the diagram is polygonal, note that any diagram is isotopic to a polygonal diagram (i.e., it has a PL structure). Now fix a point O on the plane, we can take it to be the origin on the plane, and note that if O is contained in an edge of the diagram, we can always perturb the diagram so that O is not on any edge anymore.

Now, take an edge xy on the diagram of D , obviously it inherits the orientation of the diagram. Now suppose we start at the point x on this edge, and travel along xy until we get to y . If in this travel we rotate counterclockwise in relation to the point O , we call xy a *positive* edge, on the contrary, it is a *negative* edge. Note that if all the edges are positive, it means that all the components of the link rotate around the point O and hence, up to isotopy, we already have a closed braid.

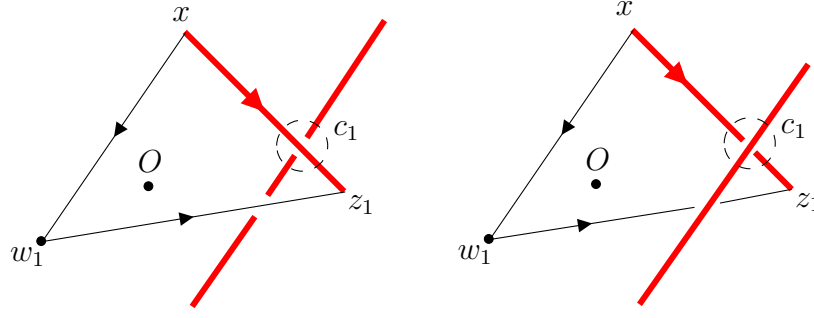
Therefore, we must prove that for every negative edge we can isotope the diagram so that the negative edge become a positive one. In fact we can do this by induction. Suppose first we take an edge and name all the crossings contained in it as $c_1, c_2, \dots, c_p$. Now, take intermediary points between these crossings, $z_1, z_2, \dots, z_{p-1}$, then z_1 is between c_1 and c_2 , z_2 is between c_2 and c_3 and so on. We will induct on the number of crossings p .

Suppose an edge have no crossings at all. Then we choose a point w on the plane in such a way that the point O is contained in the interior of the triangle xwy . Now, replace xy , by the curve $xw + wy$. By doing this we may introduce new crossings on this curve, so we must specify how the curve intersects all the crossings,



by taking all the crossings as being undercrossings (the curve $xw + wy$ passing under all the crossings) we have constructed two new positive edges only by means of triangular moves, therefore the new diagram is isotopic to the first one.

Now, suppose we have $p + 1$ crossings, and suppose by induction hypothesis, that we can isotope a negative edge with p crossings to positive edges. Taking z_1 as we have constructed, consider the edge xz_1 , and again, take a point w_1 in such a way that the point O is contained in the interior of the triangle xw_1z_1 . By replacing the edge xz_1 by the curve $xw_1 + w_1z_1$, we create new crossings, we again must specify what kinds of crossings are these. By looking at c_1 , suppose xy passes under on the crossing, then take $xw_1 + w_1z_1$ passing under all crossings, if at c_1 , xy passes over, then take $xw_1 + w_1z_1$ passing over all new crossings. Hence, we again have connected two diagrams by triangular move, hence they are again isotopic.



Note also, that since the curve $xw_1 + w_1z_1$ has two positive edges, then we reduced an edge with $p + 1$ crossings by one with p , crossings, and by using the induction hypotheses we complete the result. $\square$

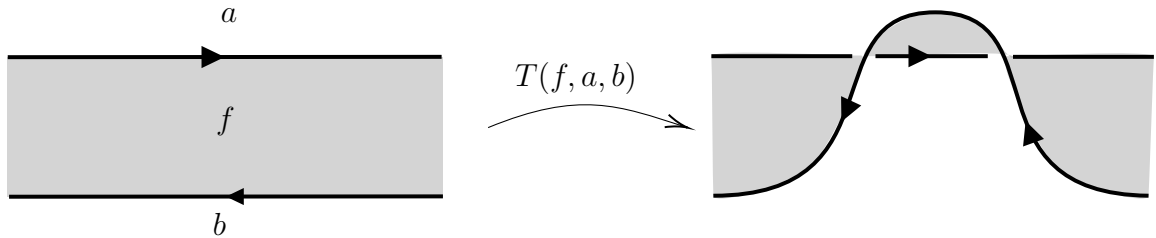
4.5 The Vogel's Algorithm

From what we've shown, we have that a given a link diagram, we can associate a closed braid to it. Also, we have given a way a of doing so. However, the given method for finding a closed braid associated to the given link can be really complicated. Thus we need a more efficient method of converting a link into a closed braid. This method is given by the Vogel's Algorithm due to P. Vogel.

First suppose we are given a diagram $D \in S^2$. We call a *face*, any connected component of $S^2 - D$. And call an *edge* any arc between two consecutive crossings on D . A triple (f, a, b) where f is a face and a and b are two edges in ∂f is said to be *admissible* if it satisfies:

- a and b are contained in different Seifert circles;
- f (hence ∂f) inherits an orientation from S^2 ; then a and b have orientations compatible to that of ∂f .

To such admissible triple (f, a, b) we can associate an R_2 move that we denote by T , i.e., $T = T(f, a, b)$ as follows:



Theorem 4.5.1. *There exists a function $\chi: \mathcal{D} \rightarrow \mathbb{Z}_+$ such that*

1. *If D is a diagram having n Seifert circles, then:*

$$n \leq \chi(D) \leq \frac{n(n+1)}{2};$$

2. *We have $\chi(D) < \chi(TD)$ if $T = T(f, a, b)$ for some admissible triple (f, a, b) ;*

3. If D is a connected diagram with n Seifert circles such that $\chi(D) \leq \frac{n(n+1)}{2}$, then there exists an admissible triple (f, a, b) in D ;
4. The connected diagram D with n Seifert circles is isotopic in S^2 to a closed braid if, and only if, $\chi(D) = \frac{n(n+1)}{2}$.

Proof. Let S be the set of Seifert circles related to $D \subset S^2$. Each oriented circle bounds two regions on S^2 , one on the left side and one on the right side when walking along the circle in the direction of the orientation.

As the n Seifert circles are disjoint, then there exists $n + 1$ connected components of the complementary of S , n of them being the disks bounded by the circles and the last one being the complement of these disks.

We can associate a graph $\Gamma(D)$ whose vertices are these connected components. We have an oriented edge from a vertex corresponding to some region to the other, if the former is the region on the left side of some Seifert circle whereas the second is the region on the right side of the same circle. We can note that by construction, this graph is actually a tree.

A m -chain c_m is the trivial graph with m edges, i.e., it simply is a segment divided into m parts.

If $c(\Gamma)$ denotes the number of oriented m -chains contained in Γ , for all $m = 1, 2, \dots$, then we set $\chi(D) = c(\Gamma(D))$ (we may, by simplicity denote $c(\Gamma(D)) = c(\Gamma)$). Counting the 1-chains we find that $\chi(D) \geq n$. Notice that $c(c_m) = \frac{m(m+1)}{2}$ since $c(m) = \sum_{i=1}^m (m-i)$.

Define a folding transformation T at the level of oriented trees as follows:

- Choose a vertex v and two edges leaving v , say vx and vy , of the oriented tree Γ ;
- The transformation T of Γ , identify two edges and two vertices, it creates one more vertex z adjacent to the vertex x (which now is the same as y), and push all edges which were before adjacent to x and y to be adjacent to z .

Geometrically, we have something like the following picture:

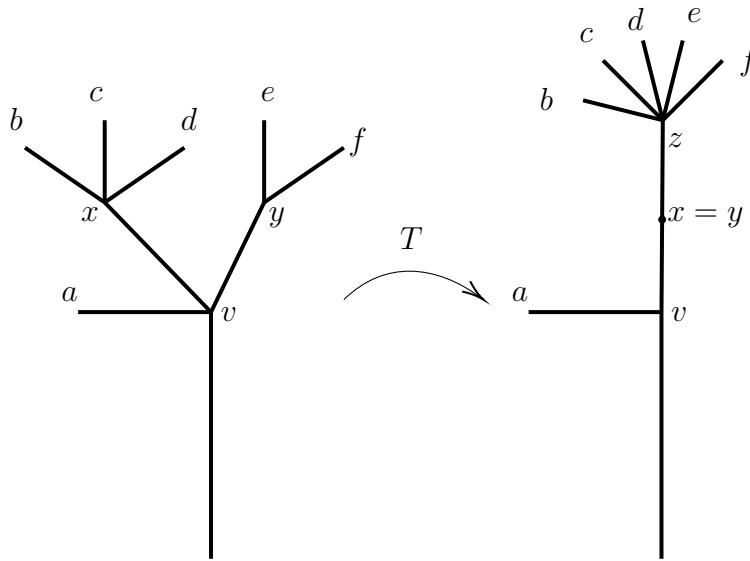
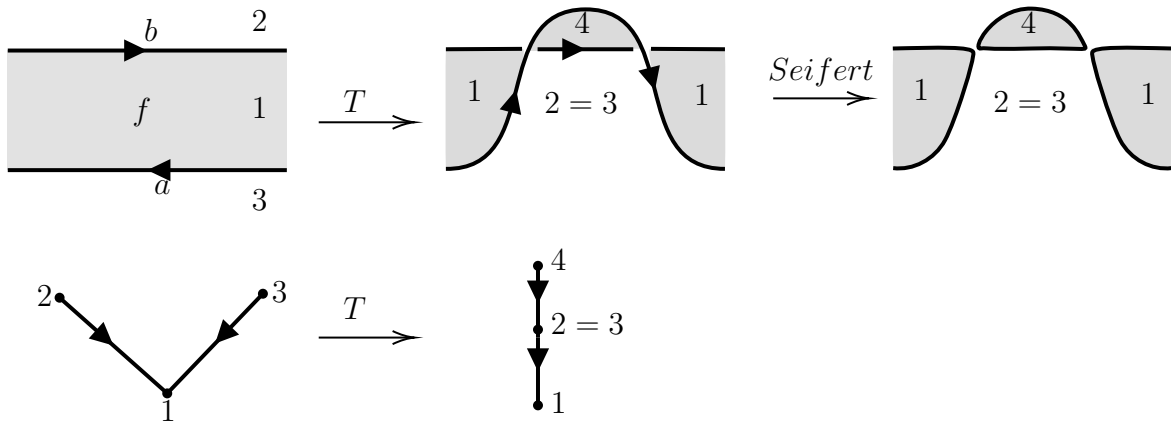


Figure 4.3: The Folding Transformation of a Tree.

Since all chains can be transported from Γ to $T\Gamma$, the chains having terminal points x or y give rise to chains having terminal points $x = y$ or z , and there is at least one more chain $xz = yz$ which does not come from Γ , thus $c(\Gamma) < c(T\Gamma)$.

We can choose this transformation as long as Γ is not a chain and, in particular, $c(\Gamma) \leq c(c_n) = \frac{n(n+1)}{2}$, whenever Γ has n edges.

Note now, that for an admissible triple (f, a, b) the transformation induced by $T(f, a, b)$ at the level of trees $\Gamma(D)$ is the transformation T associated to the vertex f and the two adjacent regions (which are the vertices for the leaving edges).

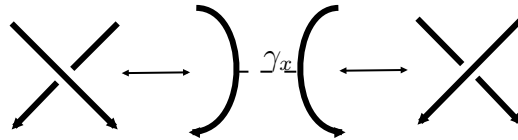


Notice also that not all T , at the level of trees, can be induced by a transformation associated to an admissible triple.

We have to prove that for every diagram D whose associated tree $\Gamma(D)$ is not a chain, should contain an admissible triple. If the tree is not a chain, there exists a vertex having at least two edges leaving it. This means that there exists two Seifert circles containing arcs from a face F , whose orientation agree with that of ∂F .

Let us say that a Seifert circle containing an arc of F is *positive* if it's orientation agree with that of ∂F (we define a *negative* Seifert circle in the same fashion). Let p be the number of positive Seifert circles and q the number of negative Seifert circles. Assume $p \geq q$ without loss of generality (we shall change the orientation of S^2 if $p < q$, then we get the same result). Also, note that $p \geq 2$.

Let's smooth every crossing to obtain a set of Seifert circles, but identify the former crossing with an small arc γ_x between the circles where the half twists should be attached.



Each arc join two Seifert circles having opposite orientations. Set $K = \bigcup_{\gamma_x \in F} \gamma_x$. Cutting open F along the arcs of K we obtain a planar domain $\hat{F} = F$, with several components, and each component corresponds to a face of the complementary of D .

Suppose there is no admissible triple (f, a, b) . Then each component of $\hat{F}$ touches exactly one positive and one negative circle.

Set f_+ for the positive circle. If $f \cap f' \neq \emptyset$ then the Seifert circles associated must be the same, i.e., $f_+ = f'_+$. Thus the function $f \rightarrow f_+$ is a locally constant function, hence it is a constant.

This means that ∂F contains only one positive Seifert circle, contradicting the fact that $p \geq 2$.

When $\Gamma(D)$ is a chain, we can isotope it on S^2 such that S is the union of concentric circles coherently oriented. Then, the arcs γ_x are transverse arcs joining consecutive circles. Replacing these arcs by crossings we obtain that D is a closed braid. $\square$

4.6 Markov's Theorem

4.6.1 Initial Concepts

We first introduce some concepts that could be introduced earlier but it was not useful until now. We define an *axis* of a link, as being the line crossing orthogonally the projection plane of the link. Note that a orientation of this axis induces an orientation of the plane in a clockwise or an anticlockwise way. In general we are going to use ℓ to denote the axis of a link. Also, in any diagram the symbol $\odot$ or $\ominus$ to indicate in which direction is the plane oriented.

First let's establish some new notation for what we are going to do from now on. We denote by $[a_1, \dots, a_n]$ the *convex hull* of the points $a_i, i = 1, \dots, n$, i.e., the smallest convex set containing all those points. We define $[a_1, \dots, a_n][b_1, \dots, b_m] \doteq [a_1, \dots, a_n] \cap [b_1, \dots, b_m]$. We say that a link is in *general position* with respect to its axis ℓ , when no edge of the link is coplanar with the axis. Some times we may denote $[a, b] = ab$ by simplicity.

Observations. In this section we work only with polygonal links, remember that any link is isotopic to a polynomial link, therefore working with polygonal links does not imply loss of generality. In the diagrams that will be presented in this section, the dashed lines represent how the link was before each movement, while the continuous lines represent how the link is after each movement.

Definition 4.6.1 (Movement of type- $\mathcal{E}$). Let V a link with side ac and b a point which is not in V . Suppose that

$$[a][b, c] = [a, b][c] = \emptyset$$

$$[a, b, c]V = [a, c]$$

then we define:

$$\mathcal{E}_{a,c}^b V = V - ac + ab + bc.$$

Thus, all the operations (movements) of the type $\mathcal{E}_{a,c}^b$ or $(\mathcal{E}_{a,c}^b)^{-1}$ are called *movements of the type $\mathcal{E}$* and we say that $\mathcal{E}$ is applicable to V .

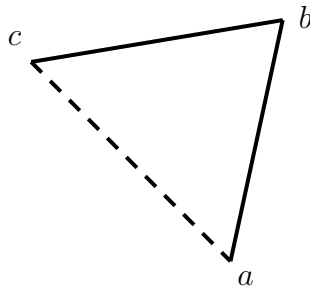


Figure 4.4: A Movement of Type- $\mathcal{E}$.

Notice that all movements of this type can create new crossings inside a diagram, let's then establish that every time we apply such movement, all new crossings created are of the same type. Since we have already introduced the triangular moves, one can easily see that the type- $\mathcal{E}$ movements are exactly the triangular moves, therefore links connected by type- $\mathcal{E}$ movements are obviously isotopic.

We say that a link is *combinatorially equivalent* to another link, if the latter can be obtained from the former by a finite sequence of applications of movement of type- $\mathcal{E}$. Since we are allowing inverses as type- $\mathcal{E}$ movements, then the combinatorial equivalence is a equivalence relation.

Now let's introduce some lemmas related to the type- $\mathcal{E}$ movement that are going to be useful for us, some of them are not going to be proved for now, we will get back to all the proves that are going to be left off now, once we conclude our main objective: proving the Markov's theorem.

Lemma 4.6.1. *Every link is combinatorially equivalent to a link in general position.*

Proof. Let V be a link that is not in general position, meaning that one of its edges is coplanar with the ℓ axis. Thus, considering ab to be such edge, we see that taking a point b' which does not lie either in V and in the plane containing ℓ and ab , we have that $[a][b, b'] = [a, b][b'] = \emptyset$ and $[a, b, b']V = [a, b]$, i.e., $\mathcal{E}_{a, b'}^b$ is applicable to V and the link $\mathcal{E}_{a, b'}^b V$ has one less coplanar edge with ℓ . Continuing like this, running through all the edges on the link, we get a new link that is isotopic to V , but it is in general position. $\square$

Note that using the Alexander's Theorem follows at once, as a corollary of the previous lemma, once that every link is combinatorially equivalent to a closed braid.

Definition 4.6.2 (Movement of type- $\mathcal{S}$). Let $a_0, \dots, a_m, b_0, \dots, b_m$ given points satisfying that:

$$[a_{i-1}, b_i] > 0, [b_i, a_i] > 0, [a_{i-1}, a_i] < 0 \text{ e } a_i \in [a_0, a_m], i = 1, \dots, m.$$

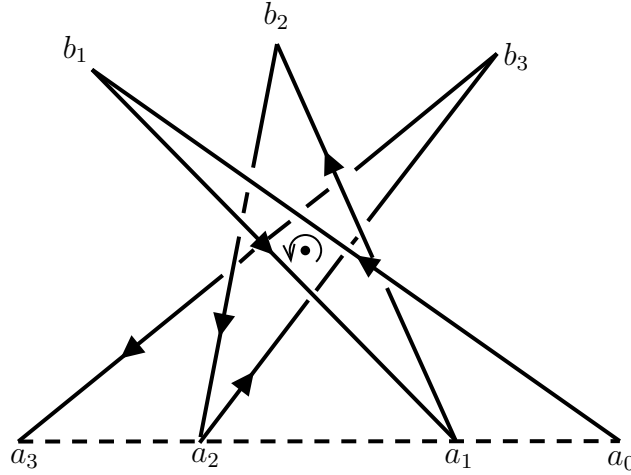
Then if

$$\sum_{i=1}^m [a_{i-1}, b_i, a_i] V = [a_0, a_m]$$

we define

$$\begin{aligned} \mathcal{S}_{a_0, \dots, a_m}^{b_1, \dots, b_m} V &= \left(\prod_{i=m}^1 \mathcal{E}_{a_{i-1}, a_i}^{b_i} \right) \left(\prod_{i=m-1}^1 \mathcal{E}_{a_{i-1}, a_m}^{a_i} \right) V \\ &= V - a_0 a_m + \sum_{i=1}^m (a_{i-1} b_i + b_i a_i) \end{aligned}$$

And we say that the operations of the type $\mathcal{S}_{a_0, \dots, a_m}^{b_1, \dots, b_m}$ and $\left(\mathcal{S}_{a_0, \dots, a_m}^{b_1, \dots, b_m} \right)^{-1}$, are of type- $\mathcal{S}$.

Figure 4.5: A Moviment of Type- $\mathcal{S}$.

The result of a type- $\mathcal{S}$ operation applied to a link, is usually called a *sawtooth*. Notice also that, we are using the definitions of positive edge and negative edge introduced previously, we are denoting a positive edge by writing $ab > 0$ and a negative side by writing $ab < 0$.

We also define an important concept called the *height* of a link, as the number of negative edges of the link. Notice that the result of applying a type- $\mathcal{S}$ operation is changing a negative edge by a sequence of positive edges, hence, we change the height of the link by exactly one unit by applying the type- $\mathcal{S}$ movement one time.

Consider the case where $m = 3$. When we apply $\mathcal{S}_{a_0, a_1, a_2, a_3}^{b_1, b_2, b_3}$ to the link V , what we are actually doing, is applying $\mathcal{E}_{a_1, a_3}^{a_2} \mathcal{E}_{a_0, a_3}^{a_1}$ to V , so that we partition the segment $[a_0, a_3]$ into new negative segments, and then, in each one of these segments, we are applying a type- $\mathcal{E}$ movement, i.e., applying $\mathcal{E}_{a_2, a_3}^{b_3} \mathcal{E}_{a_1, a_2}^{b_2} \mathcal{E}_{a_0, a_1}^{b_1}$, therefore making each one of these new negative segments two new positive segments.

Lemma 4.6.2. *Let V be a link in general position. Let $[a_0, a]$ be a negative edge V . Then a sawtooth may be erected on $[a_0, a]$. Moreover, suppose S is a simplex in the three dimensional space ($\mathbb{R}^3$ ou S^3), satisfying*

- (i) $S \cap [a_0, a] = \emptyset$ or $[a_0]$ or $[a_1]$;
- (ii) If $S \cap [a_0, a_1] \neq \emptyset$, then $S \cap [a_0, a_1]$ is a vertex of S ;
- (iii) The 1-simplices of S are not parallel to ℓ .

Then the sawtooth can be chosen to avoid S .

Proof. Let a_{i-1}, a_i be two points in a_0a such that $a_{i-1}a_i < 0$. Then the planes containing ℓ and the points a_{i-1} and a_i divide the space into 4 regions that we will denote as $I(a_{i-1}, a_i)$, $II(a_{i-1}, a_i)$, $III(a_{i-1}, a_i)$ e $IV(a_{i-1}, a_i)$, as in Figure 4.6.

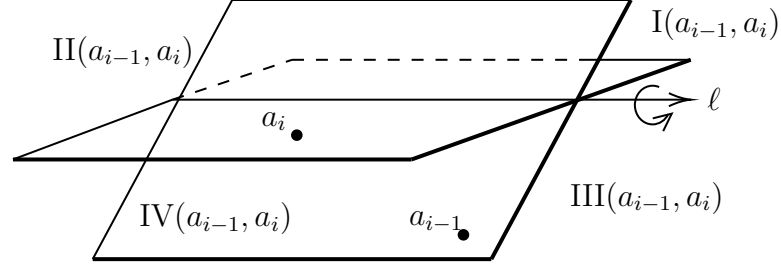


Figure 4.6: Division of space into 4 regions.

Note that if $b_i \in I(a_{i-1}, a_i)$ then we'll have that $a_{i-1}b_i > 0$, $b_i a_i > 0$. Suppose now that the projection of $a_0 a$ has no crossing (and call E_2 the projection plane, which sometimes can be thought of as being the Euclidean 2-space). Now choosing b sufficiently distant from the projection plane of V (which we shall remember, is parallel to ℓ), we can make the angle between $a_0 b a$ and the projection plane get arbitrarily next to $\frac{\pi}{2}$.

Thus, we can take $b \in I(a_{i-1}, a_i)$ in such a way that $[a_0, b, a]V = [a_0, a]$, hence $\mathcal{S}_{a_0, a}^b$ is applicable to V . More than that, $\mathcal{S}_{a_0, a}^{b_1}$ is also applicable to any point b_1 that lies above b (with respect to the projection plane of the link).

Now, since S satisfies (i), (ii) e (iii), then each triangle $a_0 b_1 a_1$ intersects S in no more than a triangle (possibly degenerate), thus, choosing b_1 even more distant from E_2 , we can erect the sawtooth in such a way that it avoids S .

In the general case, $a_0 a$ may contain a finite amount of points $p_1, \dots, p_r$ that are projected into E^2 as crossings. If that occurs, we may have that $\mathcal{S}_{a_0, a}^{b_1}$ is not applicable to V . The idea then, will be to put the teeth of the sawtooth sufficiently narrow, next to each crossing, in such a way that these teeth pass in between the edges of the link.

Let's call P_1 the plane containing ℓ and p_1 , since this plane intersects S in a convex set, we can always find a line segment beginning at p_1 and ending at ℓ avoiding both V and S . Thus, choosing a point $b_2 \in I(a_1, a_2)$ with $[a_a, a_2] \subset [a_0, a]$ we construct a triangle $[a_1, b_2, a_2]$ that avoids S and intersects V only at $[a_1, a_2]$. Thus $[a_0, a_1]$ has no crossings and we can construct a sawtooth in the same fashion we did on the previous situation, finding $b_2 \in I(a_0, a_1)$. Thus, it is enough to repeat the same construction for each p_i , $i \in \{2, \dots, r\}$, and we construct the desired sawtooth. $\square$

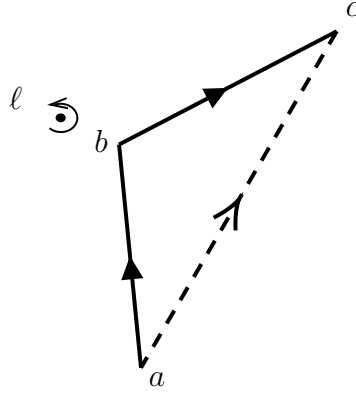
In the previous lemma, the word *avoid*, means that:

$$\left(\bigcup_{i=1}^m [a_{i-1}, b_i, a_i] \right) \cap S = \emptyset \text{ or } [a_0] \text{ or } [a_m].$$

Definition 4.6.3 (Movement of type- $\mathcal{R}$). If $\mathcal{E}_{a,c}^b$ is applicable to the link V and if $ab > 0$, $bc > 0$ and $ac > 0$, then:

$$\mathcal{R}_{a,c}^b V = \mathcal{E}_{a,c}^b V = V - ac + ab + bc.$$

Operations of the type $\mathcal{R}_{a,c}^b$ and $(\mathcal{R}_{a,c}^b)^{-1}$ are called operations of type- $\mathcal{R}$.



Note that operations of type- $\mathcal{R}$ are nothing more than type- $\mathcal{E}$ operations applied to positive edges, in such a way that the new edges created after the operation are again positive edges, therefore while type- $\mathcal{E}$ operations may change the height of a link, type- $\mathcal{R}$ does not affect the height.

Definition 4.6.4 (Movement of type- $\mathcal{W}$). Suppose $\mathcal{E}_{a,d}^b$ is applicable to V and that $\mathcal{E}_{b,d}^c$ is applicable to $\mathcal{E}_{a,d}^d V$. If

$$ab, bc, cd, ca, db, ad > 0$$

and if

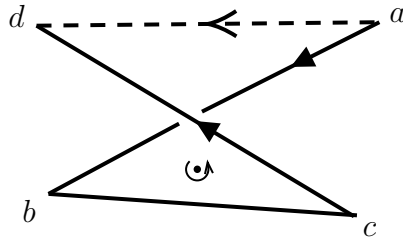
$$[a, b, c, d][V] = [a, d]$$

then we define

$$\mathcal{W}_{a,d}^{b,c} V = \mathcal{E}_{b,d}^c \mathcal{E}_{a,d}^b V = V - ad + ab + bc + cd.$$

And we say that any operation of the type $\mathcal{W}_{a,d}^{b,c}$ or $(\mathcal{W}_{a,d}^{b,c})^{-1}$ is an operation of type- $\mathcal{W}$.

Note that a type- $\mathcal{W}$ operation is a Reidemeister R_1 move, around the link axis ℓ .



4.6.2 Useful Lemmas for Proving the Markov's Theorem

Let V and V' be combinatorially equivalent closed braids. Then there exists a finite sequence of closed braids $V = V_0, V_1, \dots, V_s = V'$ such that $V_i, 1 \leq i \leq s$ is obtained from V_{i-1} by a unique application of an operation of type- $\mathcal{R}$ or type- $\mathcal{W}$, this is clear since operations of type- $\mathcal{R}$ and $\mathcal{W}$ have been defined in terms of type- $\mathcal{E}$ operations.

Definition 4.6.5. Let V and V' be links. A sequence of links

$$V = V_0 \rightarrow v_1 \rightarrow \dots \rightarrow V_n = V$$

satisfying that each link $V_i = \mathcal{F}_i V_{i-1}$, where $\mathcal{F}_i$ represents a unique application of a operation of one of the types: $\mathcal{E}, \mathcal{S}, \mathcal{R}, \mathcal{W}$ then this sequence will be called a *deformation chain* of V into V' .

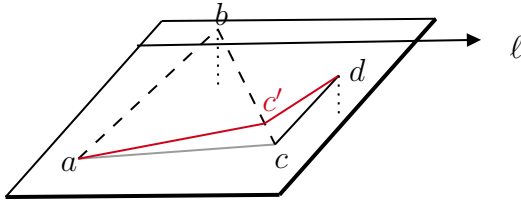
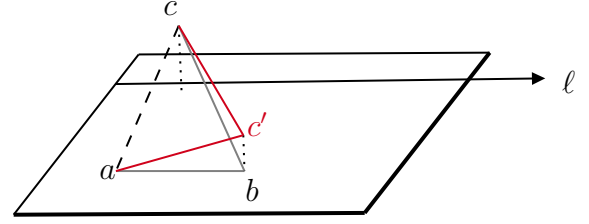
Oviously, if V e V' are combinatorially equivalent, then there exists a deformation chain between them.

We present now two lemmas that will be used to prove the Markov's Theorem, these lemmas will be proved here, since they will be directly used on the proof of Markov's Theorem.

Lemma 4.6.3. *If V and V' are in general position and are combinatorially equivalent, then there exists a deformation chain between V and V' such that all the links in the deformation chain, are in general position.*

Proof. If $\mathcal{E}_{a,c}^b$ creates at least one side that is coplanar with ℓ , then we replace $\mathcal{E}_{a,c}^b$ by $\mathcal{E}_{a,c}^{b'}$, where b' is chosen to be a point close to b , but outside the plane containing ℓ and the side which is coplanar with ℓ . After that, we replace b by b' in all the following links on the sequence.

If $(\mathcal{E}_{a,c}^b)^{-1}$ creates an edge coplanar with ℓ and if cd is the edge that is adjacent to bc , then we choose a point c' close to c and contained in bc and we replace $(\mathcal{E}_{a,c}^b)^{-1}$ by $(\mathcal{E}_{c',d}^c)^{-1} (\mathcal{E}_{c,c'}^b)^{-1} (\mathcal{E}_{b,c}^{c'})$ and replace c by c' in all the following links on the sequence.



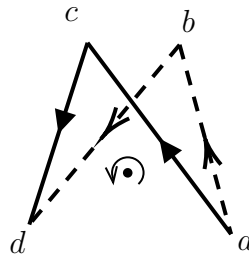
Note that this process will create two new links in the deformation chain between V and V' , however it changes in 1 unit the amount of links in this deformation chain that are not in general position. Thus, the result follows by induction in the number of links that are not in general position. $\square$

We introduce now one more type of operation, the operation of type- $\mathcal{T}$. We will also show that this new type of operation can always be expressed in terms of operations of type- $\mathcal{R}$ and $\mathcal{W}$.

Definition 4.6.6 (Operation of type- $\mathcal{T}$). Suppose that $(\mathcal{E}_{a,d}^b)^{-1}$ is applicable to V and that $\mathcal{E}_{a,d}^c$ is applicable to $(\mathcal{E}_{a,d}^b)^{-1} V$. If $ab, ac, bd, cd > 0$ and $ad < 0$ then we define:

$$\mathcal{T}_{a,d}^{b,c} V = \mathcal{E}_{a,d}^c (\mathcal{E}_{a,d}^b)^{-1} V = V - ab - bd + ac + cd.$$

All operations of the type $\mathcal{T}_{a,d}^{b,c}$ or $(\mathcal{T}_{a,d}^{b,c})^{-1}$ are called operations of type- $\mathcal{T}$.



Lemma 4.6.4. *Any operation of type- $\mathcal{T}$ can be written in terms of operations of the type $\mathcal{R}$ and $\mathcal{W}$.*

Proof. First, we show that there exists points a' and d' such that the operations $(\mathcal{R}_{c,d}^{d'})^{-1}$, $(\mathcal{W}_{d',d}^{a',b})^{-1}$, $\mathcal{W}_{a,a'}^{c,d'}$ and $\mathcal{R}_{a,b}^{a'}$ are applicable. In fact, choosing a point a' on the plane that contains $[a, b, d]$, that is close to a and in such a way so that $aa' > 0$, it follows that, since $ab > 0$, we have $a'b > 0$, hence $\mathcal{R}_{a,b}^{a'}$ is applicable to V .

Now, choosing d' a point close to d that is contained in the plane containing $[a, c, d]$ that is not the same plane that contains $[a, b, d]$, we have that $([a, b, d] + [a, c, d])V = [a, b] + [b, d]$ and since a' and d' are close to a and d respectively, we have that:

$$[a, c, d', a'] [\mathcal{R}_{a,b}^{a'}] = [a, a'].$$

Moreover, note also that $ac, cd', d'a', d'a, a'c, aa' > 0$ and then, it follows that $\mathcal{W}_{a,a'}^{c,d'}$ is applicable to $\mathcal{R}_{a,b}^{a'}V$. In a completely analogous way, it follows that the operations $(\mathcal{R}_{c,d}^{d'})^{-1}$ and $(\mathcal{W}_{d',d}^{a',b})^{-1}$ are also applicable.

Now, we show that

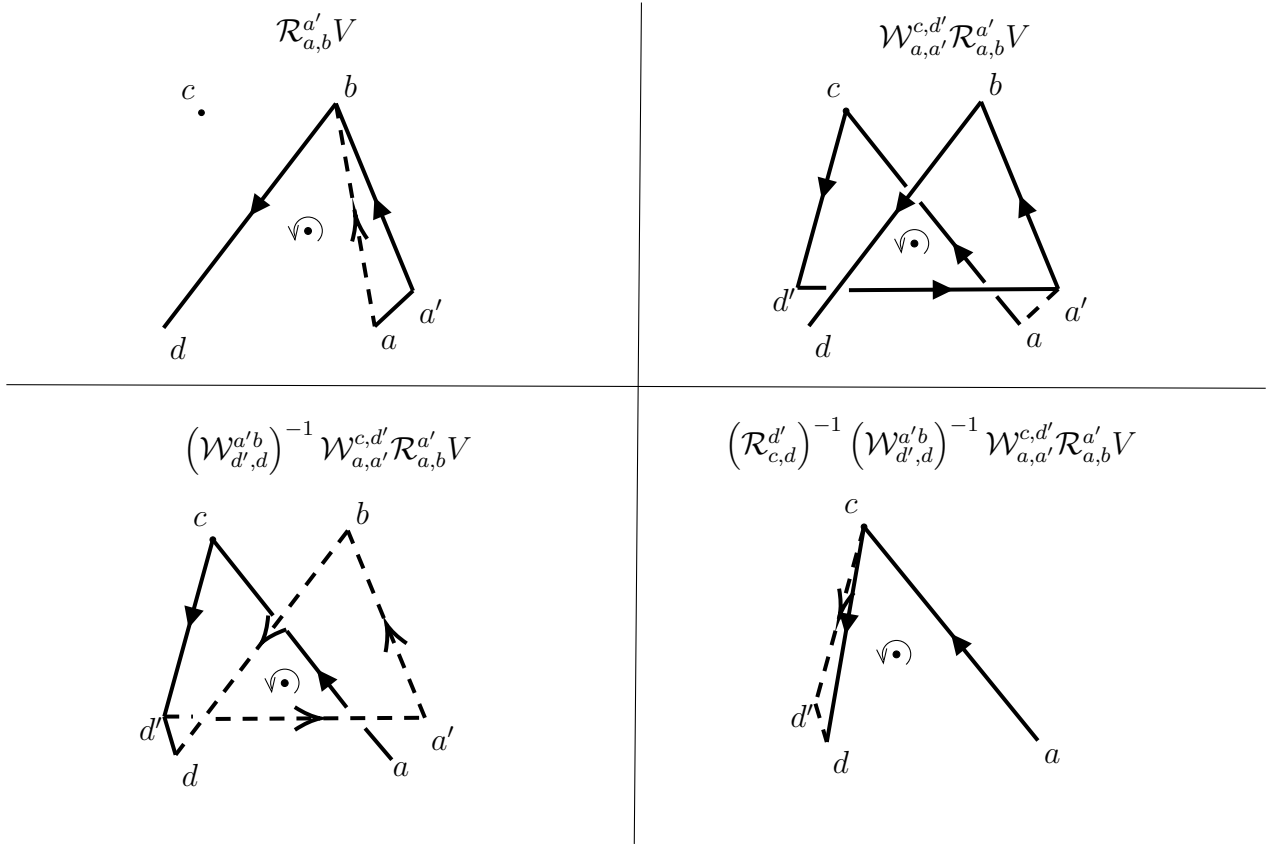
$$\mathcal{T}_{a,d}^{b,c}V = (\mathcal{R}_{c,d}^{d'})^{-1} (\mathcal{W}_{d',d}^{a',b})^{-1} \mathcal{W}_{a,a'}^{c,d'} \mathcal{R}_{a,b}^{a'} V.$$

In fact, notice that we have:

$$\begin{aligned} (\mathcal{R}_{c,d}^{d'})^{-1} (\mathcal{W}_{d',d}^{a',b})^{-1} \mathcal{W}_{a,a'}^{c,d'} \mathcal{R}_{a,b}^{a'} V &= (\mathcal{R}_{c,d}^{d'})^{-1} (\mathcal{W}_{d',d}^{a',b})^{-1} \mathcal{W}_{a,a'}^{c,d'} (V - ab + aa' + a'b) \\ &= (\mathcal{R}_{c,d}^{d'})^{-1} (\mathcal{W}_{d',d}^{a',b})^{-1} ((V - ab + aa' + a'b) - aa' + ac + cd' + d'a') \\ &= (\mathcal{R}_{c,d}^{d'})^{-1} (\mathcal{W}_{d',d}^{a',b})^{-1} (V - ab + a'b + ac + cd' + d'a') \\ &= (\mathcal{R}_{c,d}^{d'})^{-1} ((V - ab + a'b + ac + cd' + d'a') + dd' - d'a' - a'b - bd) \\ &= (\mathcal{R}_{c,d}^{d'})^{-1} (V - ab + ac + cd' + d'd - bd) \\ &= (V - ab + ac + d'd - bd) + cd - d'a' - a'b \\ &= V - ab - bd + ac + cd \doteq \mathcal{T}_{a,d}^{b,c}V \end{aligned}$$

ans that completes the proof. □

Geometrically, we can see what is going on, on the previous proof in the next figure:



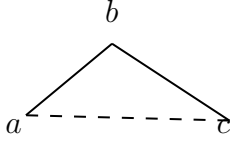
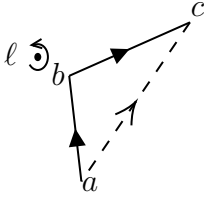
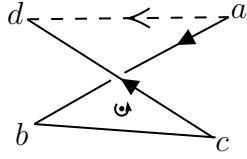
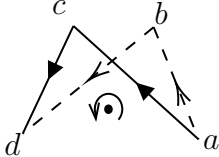
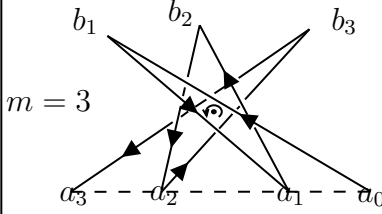
The Last Necessary Results

We begin with some important observations. First, we will be using the notation $h(V)$ to denote the height of the link V , also, remember that the height of a link is defined as the number of negative sides the link has (given an axis ℓ). Thus, we have that a link V has $h(V) = 0$ if, and only if, this link is a closed braid.

Note also that, so far we have defined the following types of moves:

- Movement of type- $\mathcal{E}$: May or may not change the link's height;
- Movement of type- $\mathcal{R}$: Does not affect the link's height;
- Movement of type- $\mathcal{W}$: Does not affect the link's height;
- Movement of type- $\mathcal{T}$: Does not affect the link's height;
- Movement of type- $\mathcal{S}$: It will always change the link's height by one unity.

The following table represents the projection of each movement and when it is applicable to the link V :

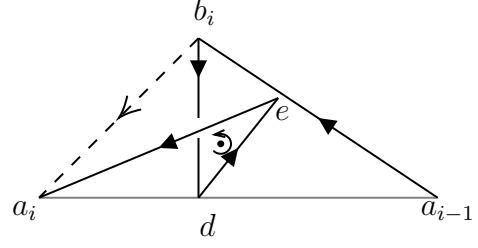
<i>Type – $\mathcal{E}$</i> 	<i>Type – $\mathcal{R}$</i> 	<i>Type – $\mathcal{W}$</i> 
$[a][b, c] = [a, b][c] = \emptyset$ $[a, b, c]V = [a, c]$	$\mathcal{E}$ – applicable $ac > 0$	$\mathcal{E}$ – applicable $ac, bc, cd, ca, db, cd > 0$
<i>Type – $\mathcal{T}$</i> 	<i>Type – $\mathcal{S}$</i> 	
$\mathcal{E}$ – applicable $ab, ac, bd, cd > 0$ $ad < 0$	$\mathcal{E}$ – applicable $a_{i-1}b_i > 0, b_ia_i > 0$ $a_{i-1}a_i < 0, a_i \in [a_0, a_m]$ $\sum_{i=1}^m [a_{i-1}, b_i, a_i][V] = [a_0, a_m]$	

Next lemma gives us an important property of the movement of type- $\mathcal{S}$, it tells us that we can transit from one sawtooth to another one, only, by means of movements of type $\mathcal{R}$ and $\mathcal{W}$, even if the sawtooths have require different partitions on the edge it is being applied.

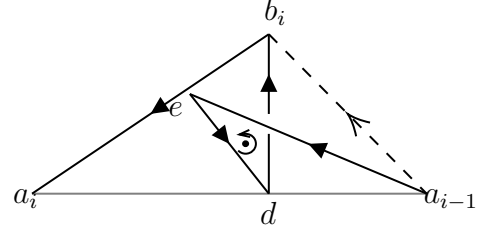
Lemma 4.6.5. $\left(\mathcal{S}_{c_0, \dots, c_m}^{d_1, \dots, d_m}\right) \left(\mathcal{S}_{a_0, \dots, a_n}^{b_1, \dots, b_n}\right)^{-1}$, where $[a_0, a_n] = [c_0, c_m]$ can be written in terms of movements of type $\mathcal{R}$ and $\mathcal{W}$.

Proof. The first step on the proof consists on showing that given a *sawtooth* in an edge of the link V , we can refine it to any point d of $[a_0, a_m]$, using only moves of type $\mathcal{W}$ and $\mathcal{T}$. We have 3 cases to consider if d is a point in $[a_{i-1}, a_i]$.

Case I: $b_i d > 0$. We choose a second point e sufficiently next to $[a_{i-1}, b_i, a_i]$ in such a way that $de > 0$ and $ea_i > 0$, thus $\mathcal{W}_{b_i, a_i}^{d, e}$ is applicable to V , and apply it to V .



Case II: $b_i d < 0$. We choose a second point e sufficiently next to $[a_{i-1}, b_i, a_i]$ in such a way that $a_{i-1}e > 0$ and $ed > 0$, thus, $\mathcal{W}_{a_{i-1}, b_i}^{e, d}$ is applicable to V , and apply it to V .



Case III: $b_i d = 0$. We apply a movement of type- $\mathcal{T}$ to choose a new vertex b'_i in such a way that $b'_i d$ follows into one of the previous cases.

The second part of the proof is to determine how to write $(\mathcal{S}_{c_0, \dots, c_m}^{d_1, \dots, d_m}) (\mathcal{S}_{a_0, \dots, a_n}^{b_1, \dots, b_n})^{-1}$ in terms of movements of the type $\mathcal{R}$ e $\mathcal{W}$. Note that we can take the segment $[a_0, a_m]$ and refine it with the points $\{a_0, \dots, a_n\} \cup \{c_0, \dots, c_m\}$. Then, using $\mathcal{T}$, if necessary, to choose adequate vertices for the sawtooth, we apply the first part of the proof and refine the sawtooth with respect to all the points.

Now, since we have defined the refinement process in terms of movements of type $\mathcal{W}$ and $\mathcal{T}$, and they are both invertible, it follows that the refinement process of a sawtooth can be undone, thus, we apply the inverse of the refinement process to the sawtooth generated by the points $a_0, \dots, a_n$ and we get the sawtooth with respect only to the points $c_0, \dots, c_m$. Thus, $(\mathcal{S}_{c_0, \dots, c_m}^{d_1, \dots, d_m}) (\mathcal{S}_{a_0, \dots, a_n}^{b_1, \dots, b_n})^{-1}$ is now written only in terms of movements of type $\mathcal{W}$ and $\mathcal{T}$, since we have proved that movements of type- $\mathcal{T}$ can be written in terms of movements of types $\mathcal{W}$ and $\mathcal{R}$, the result follows. $\square$

We introduce now the last two lemmas that will be used to prove the Markov's Theorem and, in fact, the Theorem follow in a simple way when we use these lemmas.

Lemma 4.6.6. *If $V^* \rightarrow V^{**}$ is a deformation chain and $h(V^*) = h(V^{**})$, then there exists another deformation chain*

$$V^* \rightarrow V_1 \rightarrow \dots \rightarrow V_s \rightarrow V^{**}$$

such that $s \geq 1$ and $h(V_i) < h(V^)$ for every $i \in \{1, \dots, s\}$.*

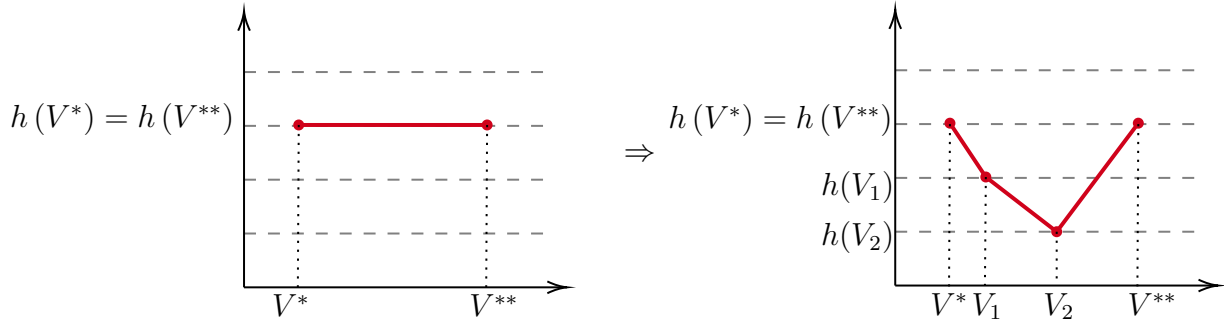
Lemma 4.6.7. *If $V^* \rightarrow \hat{V} \rightarrow V^{**}$ is a deformation chain, $h(V^*) < h(\hat{V})$ and $h(V^{**}) < h(\hat{V})$ then there exists another deformation chain*

$$V^* \rightarrow V_1 \rightarrow \dots \rightarrow V_s \rightarrow V^{**}$$

such that $s \geq 1$ and also $h(V_i) < h(\hat{V})$ for every $i \in \{1, \dots, s\}$.

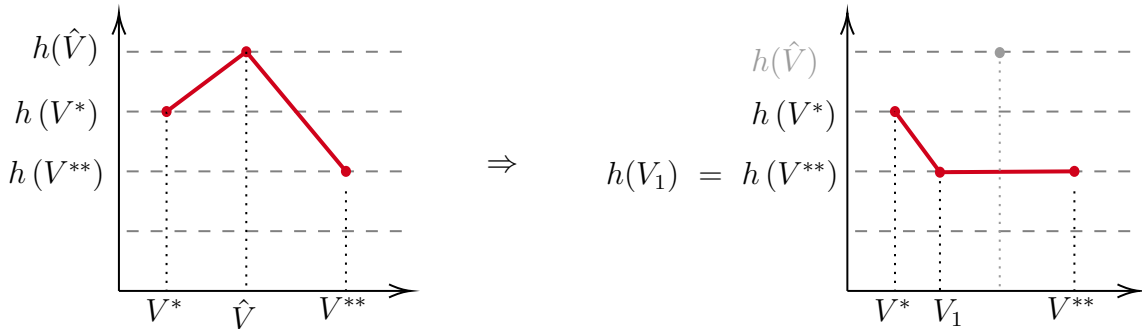
We now try to understand in a simple way what these lemmas are telling us and then, we show how the Markov's Theorem, in its first version, follows naturally. Note that if we make a graph putting the links in the horizontal axis (following their natural

order induced by their position in the deformation chain) and in the vertical axis we put their height, we have that Lemma 4.6.6, gives us the following:



Hence, if the graph of a deformation chain (constructed in the previous way) have two consecutive points with the same height, by using the lemma we can find a second path, i.e., a new graph (which is the same as saying "another deformation chain") connecting the two original points and such that all the new points on this path, have a height that is less than the the height of the original points. Observe that in this case, we are transforming a horizontal line that connects two points into a sequence of lines some of them increasing and some of them decreasing, i.e., we end up creating consecutive points with different heights.

In the case of lemma 4.6.7 we have that if two points that are connected by a path passing through one more point, and this point have a height that is greater than the height of the two edge points, then we can find a new path connecting these edge points in such a way that all the points in this new path have all heights less then the greatest height on the first, original, path. If, again, we make a graph of this two situations, we get the following



Observe that given two points with zero height, and for any path connecting this two points, wither this path have all the points with height zero, or we can apply one of the previous lemmas. This will be an important insight in proving the Markov's Theorem, since by noting that, one can see that we can always construct a path between two closed braids, made only of closed braids.

The Proof of Markov's Theorem

Theorem 4.6.1 (Markov's Theorem - First Version). *Let V and V' be two closed braids that are combinatorially equivalent. Then there exists a finite sequence of closed braids $V = V_0, V_1, V_2, \dots, V_s = V'$ such that each V_i , $i \leq 1 \leq s$, is obtained from V_{i-1} by a unique application of type- $\mathcal{R}$ or of type- $\mathcal{W}$.*

Proof. Note that by hypothesis we have that V and V' are combinatorially equivalent, then, there exists a deformation chain

$$V \rightarrow V_1 \rightarrow \cdots \rightarrow V_r \rightarrow V'$$

And by using the previous proven lemmas, we can assume that the links in this chain are all in general position, so that we can define the height of each link.

Notice now that we have two cases:

Case I: If for each $i = 1, \dots, r$, $h(V_i) = 0$, then each V_i is a closed braid, hence all operations between V_{i-1} and V_i are all of type- $\mathcal{R}$ and of type- $\mathcal{W}$, in fact, operations of type- $\mathcal{E}$ that are not operations of type- $\mathcal{R}$ and operations of type- $\mathcal{S}$ either alter the link's height, or require negative edges (i.e. strictly positive height) to be applied.

Case II: There exists i such that $h(V_i) \neq 0$. If we take $h = \max\{h(V_i); 1 \leq i \leq r\} > 0$, we have that there may exists j such that $h(V_j) = h(V_{j+1}) = j$. If that's the case, we use Lemma 4.6.6 and replace the deformation chain by the deformation chain:

$$V \rightarrow V_1^\# \rightarrow \cdots \rightarrow V_e^\# \rightarrow V'$$

where $h = \max\{h(V_i^\#); 1 \leq i \leq e\} < 0$, but there is no such j such that $h(V_j^\#) = h(V_{j+1}^\#) = h$. Hence, we are in conditions of applying Lemma 4.6.7, to construct a new deformation chain

$$V \rightarrow \tilde{V}_1 \rightarrow \cdots \rightarrow \tilde{V}_q \rightarrow V'$$

in such a way that now, $h(\tilde{V}_i) < h$ for every $i \in \{1, \dots, q\}$, thus, applying this process possibly more times we construct a deformation chain made of only closed braids. Finally, we are again in **Case I**, i.e., we have constructed a deformation chain made only of closed braids and in which, each closed braid is obtained from the previous one by only applying movements of type- $\mathcal{R}$ or of type- $\mathcal{W}$.

Finally it is important to notice the finiteness of this new deformation chain. Since the first deformation chain is finite, and each application of Lemmas 4.6.6 and 4.6.7 produces a new finite deformation chain. Now, note also that since the greatest height of a link in the deformation chain is h and h is an integer number, then by observing that an application of Lemma 4.6.7 reduces h by at least 1, then the application of the lemmas are going to be finite, hence the last deformation chain (the one made only of closed braids) is finite, which completes the proof. $\square$

Now, notice that we are missing some important proofs. In fact, the Markov's theorem follows from two lemmas that were not yet proved, we shall accomplish it now.

Proof of Lemma 4.6.6. By hypothesis, we have that $V^{**} = \mathcal{F}V^*$ where $\mathcal{F}$ is a movement of type $\mathcal{E}$, $\mathcal{S}$, $\mathcal{R}$ or $\mathcal{W}$. However, note that the movement of type $\mathcal{S}$, necessarily changes the height of the link, since we want both links in the first deformation chain, to have the same height, then $\mathcal{F} \neq \mathcal{S}$.

Also, note that the operations of type $\mathcal{E}$, may or may not alter the link height, depending on the signal of each edge created after an application of each movement, since $\mathcal{F}$ does not change the height, we look for the movements of type $\mathcal{E}$ that satisfies this condition. In fact, there exist 8 types of movements of type $\mathcal{E}$, that will be denoted by $\mathcal{E}^k$, with $k \in \{0, \dots, 7\}$, the following table represents each of these types of movements of type $\mathcal{E}$:

k	ac	ab	bc	$h(\mathcal{E}_{ac}^b V^*) - v(V^*)$	$h((\mathcal{E}_{ac}^b)^{-1} V^*) - v(V^*)$	Admissible
0	+	+	+	0	0	$\mathcal{F} = \mathcal{E}^0$
1	-	+	+	-1	+1	$\mathcal{F} \neq \mathcal{E}^1$
2	+	+	-	+1	-1	$\mathcal{F} \neq \mathcal{E}^2$
3	-	+	-	0	0	$\mathcal{F} = \mathcal{E}^3$
4	+	-	+	+1	-1	$\mathcal{F} \neq \mathcal{E}^4$
5	-	-	+	0	0	$\mathcal{F} = \mathcal{E}^5$
6	+	-	-	+2	-2	$\mathcal{F} \neq \mathcal{E}^6$
7	-	-	-	+1	-1	$\mathcal{F} \neq \mathcal{E}^7$

Thus, the only admissible moves for the lemma, are $\mathcal{E}^0$, $\mathcal{E}^3$ and $\mathcal{E}^5$. But note that $\mathcal{E}^0 = \mathcal{R}$, thus, the only movements of type $\mathcal{E}$ different from $\mathcal{R}$ and $\mathcal{W}$ are the ones of type $\mathcal{E}^3$ and $\mathcal{E}^5$. Also, note that the inverses of each movement are also admissible.

For the movements of type $\mathcal{R}$ and $\mathcal{W}$, note that since $h(V^*) > 0$, we can take a negative edge a_0a in V^* and construct a sawtooth as in Lemma 4.6.2, in such a way that this sawtooth avoids the triangle formed by $\mathcal{R}$ or the tetrahedron formed by $\mathcal{W}$. Thus, replacing $\mathcal{F}$ by $(\mathcal{S})^{-1} \mathcal{F} \mathcal{S}$, we obtain the deformation chain required by the lemma.

For the case in which $\mathcal{F} = \mathcal{E}^3$ or $\mathcal{E}^5$ and $h(V^*) > 1$, the same construction holds. Choose any other negative edge, that is not the one formed by $\mathcal{E}$ and apply the movement of type $\mathcal{S}$ as we have done before, avoiding the triangle formed by $\mathcal{E}$. Then, replace $\mathcal{F}$ by $(\mathcal{S})^{-1} \mathcal{F} \mathcal{S}$, and we obtain the deformation chain required by the lemma. The case in which $h(V^*) = 1$ must be studied with more caution, we will do that now.

The case $h(V^*) = 1$. If $\mathcal{F} = \mathcal{E}^3$, we can erect a sawtooth in ac and bc that share the vertices of the sawtooth $(b_1, \dots, b_m)$ (we say that these edges share the same sawtooth). Let us denote the base of the sawtooth on ac as a_i and the one on bc as c_i . If we assume that the edges ac and bc are sufficiently next to each other, then this sawtooth must be such that:

$$[a_i, a_{i+1}, c_i, c_{i+1}, b_i]V^* = [a_i, a_{i+1}].$$

Also, we want that, $a_i c_i \neq 0$, i.e., $a_i c_i$ is not coplanar with ℓ . We now replace $\mathcal{F}$ by a (finite) sequence of deformations in the following way:

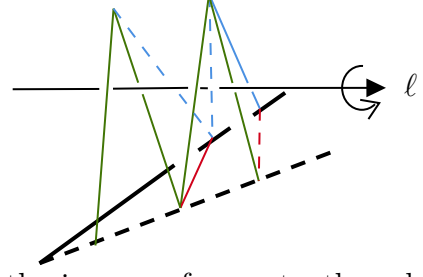
$$\mathcal{F} = (\mathcal{S}_{c_0 \dots c_m}^{b_1 \dots b_m})^{-1} \mathcal{F}_0 \dots \mathcal{F}_{m-1} \mathcal{S}_{a_0 \dots a_m}^{b_1 \dots b_m}$$

where

$$\mathcal{F}_i = \begin{cases} (\mathcal{R}_{b_i c_i}^{a_i})^{-1} \mathcal{R}_{a_i b_{i+1}}^{c_i}, & \text{se } a_i c_i > 0 \text{ e } i \neq 0 \\ (\mathcal{R}_{c_i b_{i+1}}^{a_i})^{-1} \mathcal{R}_{b_i a_i}^{c_i} & \text{se } a_i c_i < 0 \text{ e } i \neq 0 \\ \mathcal{R}_{ab_1}^b, & \text{se } i = 0 \end{cases}$$

In this sequence, the first deformation $\mathcal{S}$ erects a sawtooth on ac , reducing the height of the link. Each $\mathcal{F}_i$, can be seen as a replacement of each tooth on the sawtooth on ac to a tooth on ab that shares the same vertex,

thus, the product of all deformations of type $\mathcal{F}_i$, replaces the whole sawtooth on ac to the one on bc , except by the deformation $\mathcal{F}_0$ that creates a positive edge ab , but since it is a positive edge, it does not change the height of the link. Since each $\mathcal{F}_i$ is a product of operations of type $\mathcal{R}$, then the product does not alter the height of the link.



Finally, the deformation $\mathcal{S}^{-1}$ applied lastly, is the inverse of a sawtooth and thus, it raises the height of the link. Thus, the requirement of the lemma is satisfied for the case in which the height of the link is 1 and the edges ac and bc are sufficiently next to each other.

Considering now the general situation, we'll show that there exist points $a_{0,0}, a_{0,1}, \dots, a_{0,r}$ such that $a_{0,0} = a$, $a_{0,r} = b$ and for each pair of points $a_{0,j}, a_{0,j-1}$, we have that $a_{0,j}c$ and $a_{0,j-1}c$ admits the same construction as we have done, i.e., that these segments are sufficiently close to each other. This way, we can replace $\mathcal{F}$ for

$$\mathcal{F} = \prod_{j=1}^r (\mathcal{S}'_j)^{-1} \mathcal{F}_{0,j} \mathcal{F}_{1,j} \cdots \mathcal{F}_{m-1,j} \mathcal{S}_j$$

where $\mathcal{S}_j$ and $\mathcal{S}'_j$ are appropriate operations of type $\mathcal{S}$ and each $\mathcal{F}_{i,j}$ is a product of operations of type $\mathcal{R}$. Now, notice that, by lemma 4.6.5, we can write $(\mathcal{S})^{-1}\mathcal{S}$ as a product of operations of type $\mathcal{R}$ and $\mathcal{W}$, thus, opening the product, we have that $\mathcal{F}$ is written as $\mathcal{F} = (\mathcal{S})^{-1}\mathcal{M}\mathcal{S}$, where $\mathcal{M}$ is a (finite) product of operations of type $\mathcal{R}$ e $\mathcal{W}$, as we wanted.

We shall now show that such a sequence in fact exists. The idea will be to project V^* on the plane E^2 and study the limitations we have (based on the crossings of V^*) as we try to construct the teeth of the sawtooth in a negative edge of V^* .

- (1) Let rs be a negative edge of V^* . Note that rs has, in its projection on E^2 , at maximum a finite number of crossings. We initially consider the case where every crossing is an undercrossing of rs . Notice that if we take any point q in the region $I(rs)$, then $rq > 0$ and $qs > 0$ and even more than that, $[r, q, s]V^* = [r, s]$. Moreover, we also have that the angle formed between the triangle $[r, q, s]$ and E^2 can be taken arbitrarily close to $\frac{\pi}{2}$, by taking q sufficiently distant (above) of E^2 (since the crossings are undercrossings of rs .)
- (2) Consider rs and $r's'$ two negative, coplanar edges of V^* such that no edge of V^* lies above $[r, s, s', r']$. Let's also assume that the axis ℓ does not intersects the quadrilateral $[r, s, s', r']$. Then, it is always possible to choose a point q such that
 - (i) $rq > 0, qs > 0$;
 - (ii) $r'q > 0, qs' > 0$;
 - (iii) $[r, s, s', r', q]V^* = [r, s] \cup [r', s']$.

In fact, it suffices to notice that since ℓ does not intersects $[r, s, s', r']$, then $I(rs) \cap I(r's') \neq \emptyset$, hence, we take $q \in I(rs) \cap I(r's')$ and we have that q satisfies (i) and (ii). Now, for q to satisfy (iii) it is enough to take q sufficiently above E^2 .

- Now, if we make this line thicker in such a way that it still does not intersect the link, we can create a cone with vertex in q after crossing ℓ , furthermore, since we've taken r, s, s', r' sufficiently next to p , we have that the basis of this cone contains the quadrilateral $[r, s, s', r']$. Thus, it follows that this choice of q satisfies all the three conditions (i), (ii) and (iii). The difference between (2) and (3), is that in (3), instead of taking q distant from E_2 , we're taking q in such a way that the pyramid $[r, s, s', r', q]$ passes between the edges of V^* .

-

- (i) All edges of the link passes above the quadrilateral;
- (ii) All edges of the link passes under the quadrilateral;
- (iii) The quadrilateral is small enough and contains only one crossing corresponding to an edge passing over it and one edge passing under it.

(i) $a_{i-1,j-1}b_{i,j} > 0$, $b_{i,j}a_{i,j-1} > 0$;

- (ii) $a_{i-1,j}b_{i,j} > 0, b_{i,j}a_{i,j} > 0;$
- (iii) $[a_{i-1,j-a}, a_{i,j-1}, a_{i,j}, a_{i-1,j}]V^* = \emptyset$ or $[a_{i-1,0}, a_{i,0}]$ if $j = 0$.

Thus, for each pair $a_{0,j-1}c$ and $a_{0,j}c$, there is one shared sawtooth. The basis for these sawteeth are the points

$$a_{0,j-1}, a_{1,j-1}, \dots, a_{m,j-1} \text{ em } a_{0,j-1}c$$

$$a_{0,j}, a_{1,j}, \dots, a_{m,j} \text{ em } a_{0,j}c$$

with common vertices being $b_{1,j}, b_{2,j}, \dots, b_{m,j}$. Thus, completing the proof for the case that $\mathcal{F} = \mathcal{E}_{ac}^{b3}$.

for the case $\mathcal{F} = \mathcal{E}^5$, first consider the case in which ac and bc are sufficiently next to each other in such a way that we may erect a sawtooth in both of them, sharing the points $(d_1, \dots, d_m)$. If we denote the subdivisions in ac by $a = a_0, a_1, \dots, a_m = c$ and in bc by $b = b_0, b_1, \dots, b_m = c$. And these sawteeth should satisfy:

$$[a_i, a_{i+1}, b_i, b_{i+1}, d_i]V^* = [a_i, a_{i+1}].$$

Then, it is enough to replace $\mathcal{F}$ by

$$\mathcal{F} = \left(\mathcal{S}_{b_0 \cdot b_m}^{d_1 \dots d_m} \right)^{-1} \mathcal{F}_m \dots \mathcal{F}_2 \mathcal{F}_1 \mathcal{S}_{a_0 \dots a_m}^{d_1 \dots d_m}$$

where

$$\mathcal{F}_i = \begin{cases} \left(\mathcal{R}_{b_i d_{i+1}}^{a_i} \right)^{-1} \mathcal{R}_{a_i d_{i+1}}^{b_i} & \text{if } b_i a_i < 0 \text{ e } i \neq m \\ \left(\mathcal{R}_{d_i b_i}^{a_i} \right)^{-1} \mathcal{R}_{a_i d_{i+1}}^{b_i} & \text{if } b_i a_i > 0 \text{ e } i \neq m \\ \mathcal{R}_{d_m a_m}^{b_m} = \mathcal{R}_{d_m c}^b & \text{if } i = m \end{cases}.$$

The general case follows in the same fashion as in the case $\mathcal{E}^3$. The cases $(\mathcal{E}^3)^{-1}$ and $(\mathcal{E}^5)^{-1}$ can be obtained by using the inverses of the substitutions used in the constructions of $\mathcal{E}^3$ and $\mathcal{E}^5$, respectively. $\square$

Proof of Lemma 4.6.7. Let's write $V^{**} = \mathcal{F}_2 \mathcal{F}_1 V^*$. we shall prove the lemma for every pair $\mathcal{F}_1 \in \{\mathcal{E}^2, \mathcal{E}^4, \mathcal{E}^6, \mathcal{E}^7, \mathcal{S}^{-1}\}$ and $\mathcal{F}_2 \in \{(\mathcal{E}^2)^{-1}, (\mathcal{E}^4)^{-1}, (\mathcal{E}^6)^{-1}, (\mathcal{E}^7)^{-1}, \mathcal{S}\}$, note that $(\mathcal{E}^1)^{-1}$ may be seen as a special case of $\mathcal{S}^{-1}$.

First, note that if $\mathcal{F}_1$ increases the height of the link and $\mathcal{F}_2$ reduces it, then $\mathcal{F}_1^{-1}$ reduces the height and $\mathcal{F}_2^{-1}$ increases it, hence, if $\mathcal{F}_2 \mathcal{F}_1$ is admissible, $\mathcal{F}_2^{-1} \mathcal{F}_1^{-1}$ is also admissible. Thus, if there is a deformation chain satisfying the conclusion of the lemma for $\mathcal{F}_2 \mathcal{F}_1$, then the inverse of this chain satisfies the conclusions of the lemma for $\mathcal{F}_1^{-1} \mathcal{F}_2^{-1}$.

We then have three cases to be studied:

$$\mathcal{S} \mathcal{S}^{-1}, \mathcal{S} \mathcal{E}, \mathcal{E}^{-1} \mathcal{E}$$

since, finding a solution for $\mathcal{S} \mathcal{S}^{-1}$ implies a solution for $\mathcal{S}^{-1} \mathcal{S}$ and the same for the two other possibilities.

Case I. Suppose that $\mathcal{S} = \mathcal{S}_{c_0 \dots c_m}^{d_1 \dots d_m}$ and $\mathcal{S}^{-1} = \left(\mathcal{S}_{a_0 \dots a_n}^{b_1 \dots b_n} \right)^{-1}$. if it happens that $[a_0, a_n] = [c_0, c_m]$, from lemma 4.6.5, we have that the deformation chain may be replaced by an equivalent deformation chain written only in terms of $\mathcal{R}$ and $\mathcal{W}$, thus, since both applications do not change the link's height, we have satisfied the conclusions of the lemma.

If $[a_0, a_n] \neq [c_0, c_m]$ supposing that $h(V^*) = h$ we must have that $h(\hat{V}) = h + 1$ and $h(V^{**}) = h$. Knowing we can substitute a sawtooth in $[c_0, c_m]$ avoiding the triangle $[a_0, b_1, a_1]$, thus reducing the height from h to $h - 1$. Now, applying $(\mathcal{S}_{a_0 a_1}^{b_1})^{-1}$, increasing back the height to h . Now, substituting the sawtooth in $a_0 a_m$ by another one, now avoiding $[a_1, b_2, a_2]$, which can be done by lemma 4.6.5 only in terms of operations of type $\mathcal{R}$ and $\mathcal{W}$, and therefore not changing the link's height.

Now, in $[a_0, a_1]$ we construct a sawtooth $\mathcal{S}_{\alpha_1 \dots \alpha_k}^{\beta_1 \dots \beta_k}$ reducing the height to $h - 1$ and then applying $(\mathcal{S}_{\alpha_1 \dots \alpha_k a_2}^{\beta_1 \dots \beta_k b_2})^{-1}$, getting the height back to h .

Continuing this way in every tooth of the sawtooth in $[a_0, a_m]$, removing each one of them by a sequence of operations that do not have height greater than h . In the last application of this process, the sawtooth in $c_0 c_m$ may be put in its final position.

Case II. Suppose that $\mathcal{F} = \mathcal{SE}$ in which $\mathcal{S} = \mathcal{S}_{p_0 \dots p_m}^{q_1 \dots q_m}$ and $\mathcal{E} = \mathcal{E}_{ac}^b$. If $[p_0, p_m] \neq [b, c]$ or $[a, b]$, we can erect a sawtooth in $[p_0, p_m]$ avoiding the triangle $[a, b, c]$, say $\mathcal{S}_{c_0 \dots c_n}^{d_1 \dots d_n}$, now, substituting temporarily $\mathcal{F}$ by

$$\mathcal{S}_{p_0 \dots p_m}^{q_1 \dots q_m} (\mathcal{S}_{c_0 \dots c_n}^{d_1 \dots d_n})^{-1} \mathcal{E}_{ac}^b \mathcal{S}_{c_0 \dots c_n}^{d_1 \dots d_n}$$

thus, since $[c_0, c_n] = [p_0, p_m]$ using lemma 4.6.5 and writing $\mathcal{S}_{p_0 \dots p_m}^{q_1 \dots q_m} (\mathcal{S}_{c_0 \dots c_n}^{d_1 \dots d_n})^{-1}$ as a sequence of operations of type $\mathcal{W}$ and $\mathcal{R}$. Since $h(\mathcal{S}_{c_0 \dots c_n}^{d_1 \dots d_n} V^*) < h(V^*)$, then $h(\mathcal{E}_{ac}^b \mathcal{S}_{c_0 \dots c_n}^{d_1 \dots d_n} V^*) < h(\mathcal{E}_{ac}^b V^*) = h(\hat{V})$ and since the operations of type $\mathcal{R}$ e $\mathcal{W}$ do not change the height, the lemma is satisfied.

If $[p_0, p_m] = [b, c]$ and $\mathcal{E} \in \{\mathcal{E}^2, \mathcal{E}^6\}$ i.e., $ac > 0$, take d in ac in such a way that $dp_i > 0$ for each $i = 0, \dots, m - 1$. It is possible to take this point to d , since $cp_i > 0$ for every i , the we just take d sufficiently next to c . Thus, substituting $\mathcal{F}$ by

$$\mathcal{F} = \mathcal{S}_{p_0 \dots p_m}^{q_1 \dots q_m} \mathcal{E}_{ac}^b = (\mathcal{R}_{ab}^d)^{-1} \mathcal{W}_{dp_1}^{p_0 q_1} \dots \mathcal{W}_{dp_m}^{p_{m-1} q_m} \mathcal{R}_{ac}^b$$

each operation of type $\mathcal{W}$ is applicable since the movement $\mathcal{S}_{p_0 \dots p_m}^{q_1 \dots q_m}$ is applicable.

Now, since the operations of type $\mathcal{R}$ and $\mathcal{W}$ do not change the height, as well as it's inverses, we found a sequence that keeps constant the height in every link, therefore satisfying the lemma.

In the case where $[p_0, p_m] = [b, c]$ and $\mathcal{E} \in \{\mathcal{E}^2, \mathcal{E}^6\}$, with $ac > 0$, it suffices to take d in ac satisfying $p_i d > 0$ for every i , which is possible since $ap_i > 0$ for every i , just take d sufficiently next to a , and then, as in the previous situation, we may substitute $\mathcal{F}$ by

$$\mathcal{F} = \mathcal{S}_{p_0 \dots p_m}^{q_1 \dots q_m} \mathcal{E}_{ac}^b = (\mathcal{E}_{bc}^d)^{-1} \mathcal{W}_{p_{m-1} d}^{q_m p_m} \dots \mathcal{W}_{p_0 d}^{q_1 p_1} \mathcal{R}_{ac}^d$$

and we have the lemma satisfied.

Now, consider the situation $[p_0, p_m] = [b, c]$ and $\mathcal{E} = \mathcal{E}^7$, i.e., $ac < 0$. Consider $S = [p_0, q_1, p_1, \dots, q_m, p_m]$. Note that $[p_0, q_1, p_1] \cup \dots \cup [p_{m-1}, q_m, p_m] \subset S$. Besides S intersects ac in only one point, $c = p_m$. Thus, by lemma 4.6.2, we can erect a sawtooth in ac , avoiding S .

Suppose now that $\mathcal{S}_{e_0 \dots e_n}^{f_1 \dots f_n}$ is this sawtooth, $ac = e_o e_n$. we have that choosing d in $e_n f_n$ in such a way that $dp_i > 0$, which is possible since $bc < 0$ and then we should just take d sufficiently next to c , let's substitute:

$$\mathcal{F} = (\mathcal{E}_{ab}^c)^{-1} \left(\mathcal{S}_{e_0 \dots e_n}^{f_1 \dots f_n} \right)^{-1} \mathcal{T}_{f_n b}^{dc} \mathcal{W}_{dp_1}^{p_0 q_1} \dots \mathcal{W}_{dp_m}^{p_{m-1} q_m} \mathcal{R}_{f_n c}^d \mathcal{S}_{e_0 \dots e_n}^{f_1 \dots f_n}.$$

Note that the operations of type $\mathcal{S}$ in ac have been done avoiding S , then each operation of type $\mathcal{W}$ is applicable. Now, we know that the operation of type $\mathcal{T}$ may be written in terms of operations of type $\mathcal{R}$ and $\mathcal{W}$ by lemma 4.6.4, by doing this, we obtain a sequence satisfying the conclusions of the lemma (note that $ac < 0$, $cb > 0$ and $ab < 0$, hence, the operation $\mathcal{E}^{-1}$ that appears in this substitution, does not change the link's height).

The last situation to be considered in case II is when $[p_0, p_m] = [a, b]$ and $\mathcal{E} = \mathcal{E}^7$, i.e., $ac < 0$. Like before, let's consider the same set S and we can erect a sawtooth in ac , $\mathcal{S}_{e_0 \dots e_n}^{f_1 \dots f_n}$ avoiding S . Now, taking d sufficiently next to a and contained in $e_0 f_1$, we have that $p_i d > 0$ for every $i \in \{1, \dots, m-1\}$, thus, it is enough to substitute

$$\mathcal{F} = (\mathcal{E}_{bc}^a)^{-1} \left(\mathcal{S}_{e_0 \dots e_n}^{f_1 \dots f_n} \right)^{-1} \mathcal{T}_{cf_1}^{da} \mathcal{W}_{p_{m-1}d}^{q_m p_m} \dots \mathcal{W}_{p_0d}^{q_1 p_1} \mathcal{R}_{af_1}^d \mathcal{S}_{e_0 \dots e_n}^{f_1 \dots f_n}$$

and we conclude the lemma as before (the same observations about the movement of type $\mathcal{E}^{-1}$ that appears in this sequence is true).

Case III. $\mathcal{F} = \mathcal{E}^{-1} \mathcal{E}$, where $\mathcal{E}^{-1} = (\mathcal{E}_{df}^e)^{-1}$ and $\mathcal{E} = \mathcal{E}_{ac}^b$.

Note that $\mathcal{E}_{ac}^b$ produces at least one negative edge (note that if $k = 6$, the operation creates two negative edges). Thus, constructing $\mathcal{S}_{a_0 \dots a_m}^{b_1 \dots b_m}$ in one of these negative edges, we temporarily substitute the sequence $\mathcal{F}$ by

$$\mathcal{F} = (\mathcal{E}_{df}^e)^{-1} \left(\mathcal{S}_{a_0 \dots a_m}^{b_1 \dots b_m} \right)^{-1} \left(\mathcal{S}_{a_0 \dots a_m}^{b_1 \dots b_m} \right) (\mathcal{E}_{ac}^b).$$

Thus, this sequence is of the form $(\mathcal{E}^{-1} \mathcal{S}^{-1}) (\mathcal{S} \mathcal{E})$. Writing $\hat{V} = \mathcal{E} V^*$, $V_1 = \mathcal{S} \hat{V}$, $V_2 = \mathcal{S}^{-1} V_1$ and $V^{**} = \mathcal{E}^{-1} V_2$. We have that $h(V^*) < h(\hat{V})$ by hypothesis and $h(V_1) < h(\hat{V})$ since $\mathcal{S}$ reduces the height. Thus, $V^* \rightarrow \hat{V} \rightarrow V_1$ satisfies the conditions of case II, since it is of type $\mathcal{S} \mathcal{E}$.

In the same fashion, $V_1 \rightarrow V_2 \rightarrow V^{**}$ has $h(V_1) < h(V_2) = h(\hat{V})$ and $h(V^{**}) < h(V_2)$, which again, satisfies the conditions of case II.

Thus, it is enough to substitute the sequence $\mathcal{E}^{-1} \mathcal{E}$ by the product of the sequences obtained before and thus, we complete the proof. $\square$

4.7 Another Version of Markov's Theorem

The Markov's theorem, as we'll see, is an important tool for creating links invariants, but sometimes working with a geometric view of such theorem may be difficult, and therefore we shall study a more algebraic point of view of this theorem. We will denote by β a braid generated by the elements $\sigma_1, \sigma_2, \dots, \sigma_{n-1} \in B_n$. The notation $\hat{\beta}$ will be used to represent a closed braid, and when we wish to emphasize that the braid is constituted by n strings, we shall denote (β, n) , and we call n the *component index*.

Corollary 4.7.1. *Let $\widehat{\beta}$ and $\widehat{\beta}'$ two closed braids represented by*

$$(\beta, n) = (\sigma_{\mu_1}^{\varepsilon_1} \cdots \sigma_{\mu_t}^{\varepsilon_t}, n)$$

$$(\beta', n') = (\sigma_{\tau_1}^{\delta_1} \cdots \sigma_{\tau_k}^{\delta_k}, n').$$

Then $\widehat{\beta}$ is combinatorially equivalent to $\widehat{\beta}'$ if, and only if, there is a deformation chain

$$(\beta, n) \rightarrow (\beta_1, n_1) \rightarrow \cdots \rightarrow (\beta_s, n_s) = (\beta', n')$$

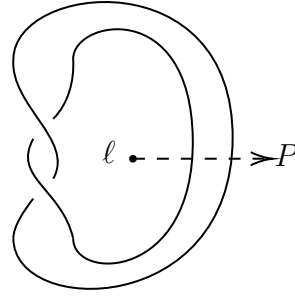
where each braid (β_{i+1}, n_{i+1}) is obtained from (β_i, n_i) by one of the following movements:

$\mathcal{M}_1$: *Replacing β_i by another word in B_{n_i} that conjugates β_i , putting $n_{i+1} = n_i$.*

$\mathcal{M}_2$: *Replacing (β_i, n_i) by $(\beta_i \sigma_{n_i}^{\pm 1}, n_i + 1)$ or, if $\beta_i = \gamma \sigma_{n_i-1}^{\pm 1}$, where γ involves only the generators $\sigma_1, \dots, \sigma_{n-2}$ replace (β_i, n_i) by $(\gamma, n_i - 1)$.*

Proof. Let $\widehat{\beta}$ and $\widehat{\beta}'$ closed braids with projections in E^2 , perpendicular to the axis ℓ .

Note that (β, n) and (β', n') can be obtained from $\widehat{\beta}$ and $\widehat{\beta}'$ by cutting the braids $\widehat{\beta}$ and $\widehat{\beta}'$ through a half plane P containing the axis ℓ , thus obtaining the braids $\beta = \sigma_{\mu_1}^{\varepsilon_1} \cdots \sigma_{\mu_t}^{\varepsilon_t}$ and $\beta' = \sigma_{\tau_1}^{\delta_1} \cdots \sigma_{\tau_k}^{\delta_k}$ from the projections of $\widehat{\beta}$ and $\widehat{\beta}'$. The integers n and n' are the cardinalities of the sets $\beta \cap P$ and $\beta' \cap P$.



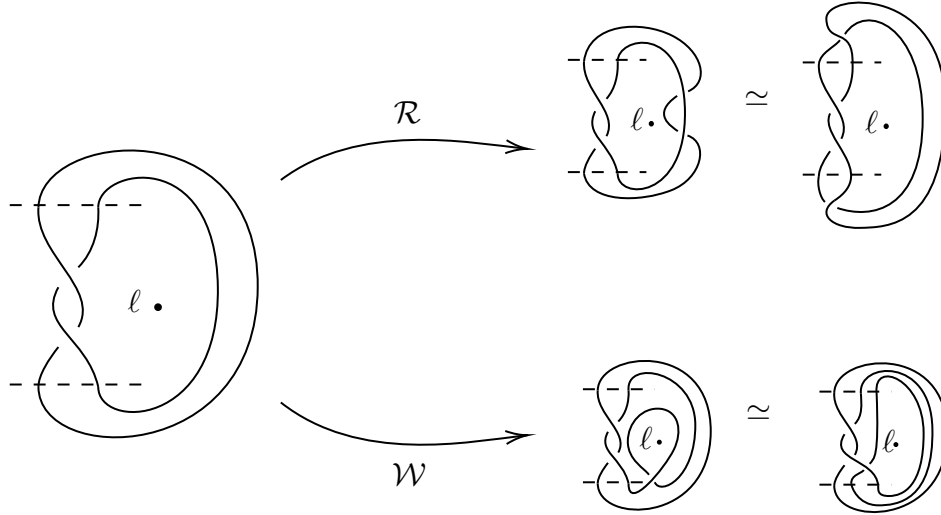
If $\widehat{\beta}$ and $\widehat{\beta}'$ are combinatorially equivalent, then by the previous version of Markov's theorem, there is a deformation chain

$$\widehat{\beta} = \beta_0 \rightarrow \widehat{\beta}_1 \rightarrow \cdots \rightarrow \widehat{\beta}_s = \widehat{\beta}'$$

with the property that each $\widehat{\beta}_i$ is a closed braid and that $\widehat{\beta}_i = \mathcal{F}_i \widehat{\beta}_{i-1}$, $q \leq i \leq s$, $\mathcal{F}_i = \mathcal{R}$ or $\mathcal{W}$.

If $\widehat{\beta}_i = \mathcal{W}_{c,d}^{a,b} \widehat{\beta}_{i-1}$, then the projection of $[a, b, c, d]$, in general, contains double points. However, we can always, by means of movements of type $\mathcal{R}$, move the rest of the link away from $[a, b, c, d]$ and after applying the movement of type $\mathcal{W}$, in such a way that $[a, b, c, d]$ will not contain any double point.

Now, to each closed braid in the deformation chain, we associate a pair (β_i, n_i) , with $\beta_i \in B_{n_{i-1}}$, with $\beta_i \in B_{n_i}$. Note that if $\beta_i = \mathcal{F}_i \beta_{i-1}$ and if $\mathcal{F}_i$ is of type $\mathcal{R}$, then the component index n_i will be the same as the index n_{i-1} . However, if $\mathcal{F}_i$ is of type $\mathcal{W}$, then $n_i = n_{i-1} \pm 1$.



Now, let $\hat{\beta}_i = (\beta_i, n_i)$ and $\hat{\beta}_i = \mathcal{R}_{a,c}^b \hat{\beta}_{i-1}$. Since the operations of type $\mathcal{R}$ preserve the component indexes, it follows that $n_i = n_{i-1}$ and that β_i and β_{i-1} represent the conjugate elements of B_n , hence, β_i is obtained from β_{i-1} by a movement of type $\mathcal{M}_1$.

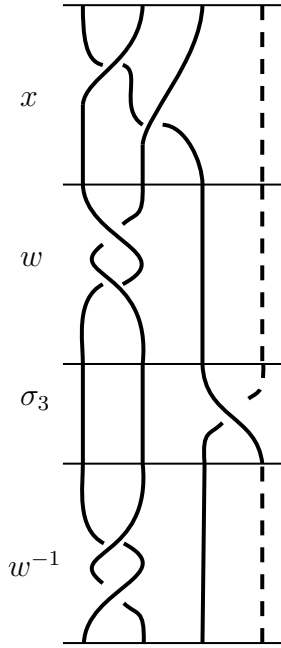
If $\hat{\beta}_i = (\mathcal{W}_{c,d}^{a,b})^{\pm 1} \beta_{i-1}$, then, as we have seen, $n_i = n_{i-1} \pm 1$. Since $[a, b, c, d]$ does not contain any double points, $(\mathcal{W}_{c,d}^{a,b})^{\pm 1}$ may be seen as a unique application of a movement of type $\mathcal{M}_2$.

For the other direction, since there is a deformation chain between the braids, it follows immediately that they are combinatorially equivalent. $\square$

In the previous corollary, we have established two important moves that can be performed in a link which we called $\mathcal{M}_1$ and $\mathcal{M}_2$, however, these moves have important names that comes from an algebraic intuition, in fact we call the movement $\mathcal{M}_1$ a *conjugation* and $\mathcal{M}_2$ a *stabilization*, or a *destabilization* if it is the inverse of a movement $\mathcal{M}_2$. With that in mind, we can rewrite the Markov's theorem as follows:

Theorem 4.7.1 (Markov's Theorem). *Closed braids are equivalent (as links) if, and only if, we can relate them by a sequence in which each term is a conjugation or a stabilization (or even a destabilization).*

The Improved Markov's Theorem



First let's define a new movement between braids called *generalized stabilization*. Let $x \in B_n$ be a braid in n components, we define a generalized stabilization of x as being the braid in B_{n+1} $xw^{-1}\sigma_n^{\pm 1}w \in B_{n+1}$ where $w \in B_n$.

Note that the generalized stabilization is like doing the product of the braid x with a stabilization movement conjugate by the braid w . Note also that if $x_1 \rightarrow x_2$ is a generalized destabilization, i.e., the inverse of a generalized stabilization, then $x_2 \rightarrow x_1$ is a generalized stabilization.

Lemma 4.7.1. *Let $x \rightarrow x'$ be a generalized stabilization, then $x \rightarrow x'$ can be written as $x \rightarrow y \rightarrow z \rightarrow x'$, with $x \rightarrow y$ a conjugation, $y \rightarrow z$ a stabilization and $z \rightarrow x'$ a conjugation.*

Proof. By hypothesis we have:

$$x' = xw\sigma_n^{\pm 1}w^{-1}.$$

Now taking $y = wxw^{-1}$ and $z = wxw^{-1}\sigma_n^{\pm 1}$ we have that the first application is a conjugation by w , the second, a stabilization and finally, a conjugation by w^{-1} , obtaining thus the sequence we wanted. $\square$

Lemma 4.7.2. *Let $x_0 \rightarrow x_1$ conjugation and $x_1 \rightarrow x_2$ generalized stabilization, then there exists y such that $x_0 \rightarrow y$ is a generalized stabilization and $y \rightarrow x_2$ is a conjugation.*

Proof. By hypothesis, we have that $x_1 = zx_0z^{-1}$ and $x_2 = (zx_0z^{-1})w\sigma_n^{\pm 1}w^{-1}$. Take then $y = x_0(z^{-1}w)\sigma_n^{\pm 1}w^{-1}z$ and now, applying the conjugation by z , $y \rightarrow x_2$, we obtain:

$$zx_0z^{-1}w\sigma_n^{\pm 1}w^{-1}zz^{-1} = zx_0z^{-1}w\sigma_n^{\pm 1}w^{-1} = x_2.$$

$\square$

Lemma 4.7.3. *Let $x_0 \rightarrow x_1$ in the generalized stabilization and $x_1 \rightarrow x_2$ a generalized stabilization. Then, there are y and z such that $x_0 \rightarrow y$ is a generalized stabilization, $y \rightarrow z$ a conjugation and $z \rightarrow x_2$ a generalized destabilization.*

Proof. We have that if $\varepsilon, \delta \in \{-1, 1\}$, then by hypothesis

$$x_0 = x_1w^{-1}\sigma_n^{\varepsilon}w$$

$$x_2 = x_1t^{-1}\sigma_n^{\delta}t$$

with $x_0, x_w \in B_{n+1}$ e $x_1, w, t \in B_n$. Let's then write $y = x_0(\sigma_n t)^{-1} \sigma_n^\delta(\sigma_n t)$ in the generalized stabilization of x_0 and $z = \sigma_{n+1} y \sigma_{n+1}$ in the conjugation of y . That is,

$$\begin{aligned} z &= \sigma_{n+1}^{-1} \left(x_0(\sigma_n t)^{-1} \sigma_n^\delta(\sigma_n t) \right) \sigma_{n+1} \\ &= \sigma_{n+1}^{-1} \left(x_1 w^{-1} \sigma_n^\varepsilon w (\sigma_n t)^{-1} \right) \sigma_n^\delta \sigma_n t \sigma_{n+1} \\ &= \sigma_{n+1}^{-1} x_1 w^{-1} \sigma_n^\varepsilon w t^{-1} \sigma_n^{-1} \sigma_n^\delta \sigma_n t \sigma_{n+1} \end{aligned}$$

and from the properties that define the braid groups, we have that $\sigma_n \sigma_{n+1} \sigma_n = \sigma_{n+1} \sigma_n \sigma_{n+1}$ also, since $x_1 w \in B_n$ and $\sigma_{n+1} \in B_{n+2}$, we have that $x_1 w^{-1}$ and σ_{n+1} commute, as well as σ_{n+1} and t commute, it follows that

$$\begin{aligned} z &= \sigma_{n+1} x_1 w^{-1} \sigma_n^\varepsilon w t^{-1} \sigma_{n+1} \sigma_n^\delta \sigma_{n+1}^{-1} t \sigma_{n+1} \\ &= x_1 w^{-1} \sigma_{n+1}^{-1} \sigma_n^\varepsilon w t^{-1} \sigma_{n+1} \sigma_n^\delta t \\ &= x_1 w^{-1} \sigma_{n+1} \sigma_n^\delta \sigma_{n+1} w t^{-1} \sigma_n^\delta t \\ &= x_1 w^{-1} \sigma_n \sigma_{n+1}^\varepsilon \sigma_n w t^{-1} \sigma_n^\delta t \\ &= x_1 t^{-1} \sigma_n^\delta t (\sigma_n^{-1} w t^{-1} \sigma_n^\delta t)^{-1} \sigma_{n+1}^\varepsilon (\sigma_n^{-1} w t^{-1} \sigma_n^\delta t) \end{aligned}$$

which proves the result, since now z is a generalized stabilization of x_2 , as we wish. $\square$

Together, these results show that:

Theorem 4.7.2 (Improved Markov's Theorem). *If the closed braids $\hat{x}$ and $\hat{y}$ are combinatorially equivalent, there exists a sequence of movements*

$$x = x_0 \rightarrow x_1 \rightarrow \cdots \rightarrow x_k \rightarrow y_k \rightarrow y_{k-1} \rightarrow \cdots \rightarrow y_0 = y$$

such that for each j , we have

- $x_j \rightarrow x_{j+1}$ is a conjugation or a stabilization.
- $y_{j+1} \rightarrow y_j$ is a conjugation or a destabilization.

In fact, lemma 4.7.3 assures us that the sequence between the braids can be written as a sequence of generalized stabilizations followed by a sequence of generalized destabilizations. Then, lemma 4.7.1 ensures us that each generalized stabilization (and similarly each generalized destabilization), may be written as a conjugation followed by a stabilization (similarly destabilization) followed by a conjugation.

5 Quantum Link Invariants

5.1 Markov's Traces

Let's consider $\mathbb{C}B_n$ as the group algebra of B_n with coefficients in the field $\mathbb{C}$, i.e., $\mathbb{C}B_n$ is the set of all linear combinations $\sum_{x \in B_n} z_x x$ with $z_x \in \mathbb{C}$, together with the field operations of $\mathbb{C}$ and with the group operations of B_n .

Thus, note that the embedding $B_n \hookrightarrow B_{n+1}$ that sends each generator σ_i into σ_i induces an embedding $\mathbb{C}B_n \hookrightarrow \mathbb{C}B_{n+1}$. Now, if we consider the family of linear functionals $t_n: \mathbb{C}B_n \rightarrow \mathbb{C}$ satisfying:

$$t_n(xy) = t_n(yx) \quad \forall x, y \in \mathbb{C}B_n$$

$$t_{n+1}(x\sigma_n) = z t_n(x) \quad \forall x \in \mathbb{C}B_n \quad \text{and} \quad t_{n+1}(x\sigma_n^{-1}) = \bar{z} t_n(x) \quad \forall x \in \mathbb{C}B_n$$

for some $z \in \mathbb{C}^*$ (here, $\bar{z}$ is the complex conjugate of z), then we call this family of linear functionals a *Markov's trace*.

Definition 5.1.1. Let $x \in B_n$, written as $\prod_{i=1}^{n-1} \sigma_{j_i}^{\lambda_i}$, we define the *exponent sum* of x the value

$$e(x) = \sum_{i=1}^{n-1} \lambda_i.$$

The following proposition shows us how we can obtain a link invariant using the Markov theorem studied previously.

Proposition 5.1.1. *If $\{t_n\}$ is a Markov's trace, the function F that associates to each closed braid $\hat{x}$, $x \in B_n$, the value:*

$$F(\hat{x}) = z^{-\frac{e(x)+n}{2}} \bar{z}^{\frac{e(x)-n}{2}} t_n(x)$$

is a link invariant, with $|z| = 1$.

Proof. to check that F is in fact a link invariant, it follows from the Markov's theorem that it is enough to check that it is invariant with respect to conjugations and stabilizations, we shall do this now.

Conjugation. Let's compute it by parts. Note that if $c \in B_n$ is a braid, then $c = \prod_{i=1}^{n-1} \sigma_{k_i}^{\alpha_i}$, which implies that $c^{-1} = \prod_{i=1}^{n-1} \sigma_{k_i}^{-\alpha_i}$, thus, we must have that by writing $x = \prod_{i=1}^{n-1} \sigma_{j_i}^{\lambda_i}$, then

$$cxc^{-1} = \left(\prod_{i=1}^{n-1} \sigma_{k_i}^{\alpha_i} \right) \left(\prod_{i=1}^{n-1} \sigma_{j_i}^{\lambda_i} \right) \left(\prod_{i=1}^{n-1} \sigma_{k_i}^{-\alpha_i} \right)$$

in other words,

$$e(cxc^{-1}) = \sum_{i=1}^{n-1} (\alpha_i + \lambda_i - \alpha_i) = \sum_{i=1}^{n-1} \lambda_i = e(x).$$

Now, since t_n is a Markov's trace, we have that $t_n(cxc^{-1}) = t_n(cc^{-1}x) = t_n(x)$, i.e., putting together this information, we get:

$$F(cxc^{-1}) = z^{-\frac{e(cxc^{-1})+n}{2}} \bar{z}^{\frac{e(cxc^{-1})-n}{2}} t_n(cxc^{-1}) = z^{-\frac{e(x)+n}{2}} \bar{z}^{\frac{e(x)-n}{2}} t_n(x) = F(x).$$

Stabilization. We proceed in the same fashion as we did for the conjugation. Since a stabilization involves a product on the right by the element $\sigma_n^{\pm 1} \in B_{n+1}$, then it's exponent is $+1$ or -1 , let's assume it is $+1$. Then, we have that $e(x\sigma_n) = e(x) + 1$. Now, since t_n is a Markov's trace, we have that $t_{n+1}(x\sigma_n) = z t_n(x)$. Now, we have that

$$\begin{aligned} F(x\sigma_n) &= z^{-\frac{e(x\sigma_n)+n}{2}} \bar{z}^{\frac{e(x\sigma_n)-n}{2}} t_{n+1}(x) \\ &= z^{-\frac{e(x)+n}{2}} z^{-\frac{1}{2}} \bar{z}^{-\frac{1}{2}} \bar{z}^{\frac{e(x)-n}{2}} \bar{z}^{\frac{1}{2}} z t_n(x) \\ &= F(x) z^{\frac{1}{2}} \bar{z}^{\frac{1}{2}} = F(x)(z\bar{z})^{\frac{1}{2}} = F(x)|z| = F(z) \end{aligned}$$

□

5.2 The Burau Representation

Definition 5.2.1. Let G be a finite group and V a vector space over a field $\mathbb{K}$, we define a *linear representation of G* or just *representation*, a homomorphism from G to $\text{GL}_n(\mathbb{K})$.

We recall that $\text{GL}(\mathbb{K})$ is the group formed by the invertible matrices (with non zero determinant) with entries on the field $\mathbb{K}$, together with the traditional operations of matrices.

The Burau representation is a linear representation of the braid group B_n , in the group $\text{GL}(\Lambda)$, where $\Lambda = \mathbb{Z}[t^{-1}, t]$, is the ring of Laurent polynomials. Such a representations, sends a generator $\sigma_i \in B_n$ in the matrix

$$I_{i-1} \oplus \begin{pmatrix} 1-t & t \\ 1 & 0 \end{pmatrix} \oplus I_{n-i-1}.$$

Where I_k is the identity matrix $k \times k$ and the non-trivial block 2×2 appears on the i -th e $(i+1)$ -th rows and columns of the matrix.

Example If we consider the generator $\sigma_4 \in B_8$, then it's image in $\text{GL}(\Lambda)$ is the matrix:

$$\begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1-t & t & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

Thus, taking a braid $x \in B_n$ which is written as a a product of powers of the generators in B_n , we can represent it as a matrix by taking its image through the representation homomorphism.

5.3 Reduced Burau Representation

Recall that $\Lambda = \mathbb{Z}[t, t^{-1}]$, then we have that the Burau representation is a map $\rho: B_n \rightarrow \text{GL}_n(\Lambda)$ such that

$$\rho(\sigma_i) = I_{i-1} \oplus \begin{pmatrix} 1-t & t \\ 1 & 0 \end{pmatrix} \oplus I_{n-i-1}.$$

The Burau representation is reducible, in the sense that it can be written as a direct sum of a $(n-1)$ -dimensional representation, which we call the *reduced Burau representation* and a 1-dimensional trivial representation.

Let x be a braid, written as $x = \prod_i \sigma_{j_i}^{\lambda_i}$. We know that

$$\rho(x) = \rho\left(\prod_i \sigma_{j_i}^{\lambda_i}\right) = \prod_i \rho(\sigma_{j_i})^{\lambda_i}$$

thus, to find the reduced Burau representation, we just conjugate $\rho(\sigma_{j_i})$ by the matrix:

$$\begin{pmatrix} 1 & 0 & 0 & \cdots & 0 \\ 1 & t & 0 & \cdots & 0 \\ 1 & t & t^2 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & t & t^2 & \cdots & t^{n-1} \end{pmatrix}$$

and remove from the resultant matrix the n -th row and column.

Example. In B_5 , we shall compute the reduced Burau representation of σ_3 . Let's call $\rho_B: B_5 \rightarrow \text{GL}(\Lambda)$ the reduced representation and $\rho: B_5 \rightarrow \text{GL}(\Lambda)$ the Burau representation. We know that:

$$\rho(\sigma_3) = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & t-1 & t & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

Now, we also have that

$$\begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 1 & t & 0 & 0 & 0 \\ 1 & t & t^2 & 0 & 0 \\ 1 & t & t^2 & t^3 & 0 \\ 1 & t & t^2 & t^3 & t^4 \end{pmatrix}^{-1} = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ -\frac{1}{t} & \frac{1}{t} & 0 & 0 & 0 \\ 0 & -\frac{1}{t^2} & \frac{1}{t^2} & 0 & 0 \\ 0 & 0 & -\frac{1}{t^3} & \frac{1}{t^3} & 0 \\ 0 & 0 & 0 & -\frac{1}{t^4} & \frac{1}{t^4} \end{pmatrix}$$

hence,

$$\begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 1 & t & 0 & 0 & 0 \\ 1 & t & t^2 & 0 & 0 \\ 1 & t & t^2 & t^3 & 0 \\ 1 & t & t^2 & t^3 & t^4 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & t-1 & t & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ -\frac{1}{t} & \frac{1}{t} & 0 & 0 & 0 \\ 0 & -\frac{1}{t^2} & \frac{1}{t^2} & 0 & 0 \\ 0 & 0 & -\frac{1}{t^3} & \frac{1}{t^3} & 0 \\ 0 & 0 & 0 & -\frac{1}{t^4} & \frac{1}{t^4} \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & t & -t & 1 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

$$\Rightarrow \rho_B(\sigma_3) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & t & -t & 1 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

Thus we can write a more general form of the reduced Burau representation as:

$$\rho_B(\sigma_1) = \begin{pmatrix} -t & 1 \\ 0 & 1 \end{pmatrix} \oplus I_{n-3} \text{ and } \rho_B(\sigma_j) = I_{j-2} \oplus \begin{pmatrix} 1 & 0 & 0 \\ t & -t & 1 \\ 0 & 0 & 1 \end{pmatrix} \oplus I_{n-j-1}$$

$$\rho_B(\sigma_{n-1}) = I_{n-3} \oplus \begin{pmatrix} 1 & t \\ 0 & -t \end{pmatrix}$$

Example. By using the previous expression, let's write $\rho_B(\sigma_4)$ in B_8 :

$$\rho_B(\sigma_4) = I_2 \oplus \begin{pmatrix} 1 & 0 & 0 \\ t & -4 & 1 \\ 0 & 0 & 1 \end{pmatrix} \oplus I_2 = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & t & -t & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}.$$

Definition 5.3.1. We say that a representation $\rho: G \rightarrow \text{GL}_n(V)$ of a group is *faithful*, when ρ is injective.

It is proven that for $n \leq 3$, the Burau representation is faithful, while for $n \geq 5$, it isn't. The case $n = 4$ is still an open question.

Theorem 5.3.1. If a link L is the closure of a braid $x \in B_n$, then

$$\widetilde{\nabla}_L(t) = \frac{(1-t)}{(1-t^n)} \det(\rho_B(x) - I_{n-1})$$

is a invariant of L , i.e., it does not depend on the braid x of which L is it's closure. $\widetilde{\nabla}_L$ is called reduced Alexander polynomial.

A proof for this theorem can be found in *Braids, Links and Mapping Class Groups* by Joan S. Birman.

Example. Let us use what has been done before, to compute the reduced Burau representation and the reduced Alexander polynomial for the figure eight knot. A braid which it's closure is given by this knot is $x = \sigma_1 \sigma_2^{-1} \sigma_1 \sigma_2^{-1}$ (represented on Figure 5.1).

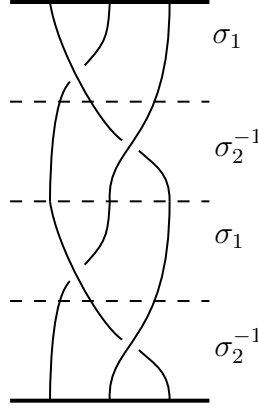


Figure 5.1: A braid that represents the figure eight knot.

Now we have that

$$\rho_B(\sigma_1) = \begin{pmatrix} -t & 1 \\ 0 & 1 \end{pmatrix} \text{ e } \rho_B(\sigma_2) = \begin{pmatrix} 1 & t \\ 0 & -t \end{pmatrix}$$

since $\rho_B(\sigma_2^{-1}) = [\rho_B(\sigma_2)]^{-1}$, we have

$$\rho_B(\sigma_2^{-1}) = \begin{pmatrix} 1 & 1 \\ 0 & -t^{-1} \end{pmatrix}$$

hence, we may write:

$$\begin{aligned} \rho_B(x) &= \rho_B(\sigma_1 \sigma_2^{-1} \sigma_1 \sigma_2^{-1}) = \rho_B(\sigma_1) [\rho_B(\sigma_2)]^{-1} \rho_B(\sigma_1) [\rho_B(\sigma_2)]^{-1} \\ &= \begin{pmatrix} -t & 1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 0 & -t^{-1} \end{pmatrix} \begin{pmatrix} -t & 1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 0 & -t^{-1} \end{pmatrix} \\ &= \begin{pmatrix} t^2 & \frac{t^4+2t^2+1}{t^2} \\ 0 & \frac{1}{t^2} \end{pmatrix} \end{aligned}$$

which is the reduced Burau representation for the figure eight knot.

Now, we have

$$\begin{aligned} &\begin{pmatrix} t^2 & \frac{t^4+2t^2+1}{t^2} \\ 0 & \frac{1}{t^2} \end{pmatrix} - \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} t^2-1 & \frac{t^4+2t^2+1}{t^2} \\ 0 & \frac{1-t^2}{t^2} \end{pmatrix} \\ \Rightarrow \det \begin{pmatrix} t^2-1 & \frac{t^4+2t^2+1}{t^2} \\ 0 & \frac{1-t^2}{t^2} \end{pmatrix} &= \frac{(-t^2+1)(t^2-1)}{t^2} = -t^2+2-\frac{1}{t^2} \end{aligned}$$

Therefore,

$$\widetilde{\nabla}_L(t) = \frac{1-t}{(1-t^3)} \left(-t^2+2-\frac{1}{t^2} \right) = \frac{-t^4+2t^2-1}{t^4+t^3+t^2}$$

5.4 Some Consequences of Studying Representations of B_n

We now introduce two main algebras that arise from the study of representations of the braid group B_n . Our focus now will be only to expose some algebraic consequences of such a study and to see how powerful such a study can be, in terms of creating brand new theories to possibly be used in other areas of mathematics, and other areas of science. We will not focus on proofs or establishing too many results, since this would go on a different direction of the main goal of this dissertation.

5.4.1 Hecke Algebras

In this subsection, and the next one, where we talk about the BMW-algebras, we do introduce the concepts, but do not go much further into details since this would deviate a little bit of the main subject of this dissertation. The main goal here is only to introduce some objects which arise from the study of representations of the braid group B_n . As we'll see, in a sense, they generalize the notion of a permutation (as in S_n) by not allowing that $b_i^2 = \text{Id}$, which is in convergence with the fact that in a braid group the same relation ($b_i^2 = \text{Id}$) is also not valid.

We start by remembering the definition of an algebra over a field $\mathbb{K}$.

Definition 5.4.1. An algebra A over a field $\mathbb{K}$ (or a $\mathbb{K}$ -algebra) is a $\mathbb{K}$ -vector space together with a multiplication operation (usually denoted by $\cdot$), satisfying:

- (a) $x \cdot (y + z) = x \cdot y + x \cdot z$, for $x, y, z \in A$;
- (b) $(y + z) \cdot x = y \cdot x + z \cdot x$, for $x, y, z \in A$;
- (c) If $\alpha, \beta \in \mathbb{K}$, then $(\alpha x \cdot \beta y) = (\alpha \beta) x \cdot y$, where $x, y \in A$.

This means that, since every algebra is a vector space and every vector space is a module over a field, then every algebra can be seen also as a module, hence we can talk about a basis for an algebra.

Let $q \in \mathbb{C}$ be any non zero complex number, we define the *Hecke Algebra* $H(q, n)$ as the quotient of the group algebra $\mathbb{C}[B_n]$ with the relation (called *quadratic relation*):

$$\sigma_i^2 = (1 - q)\sigma_i + q \text{ for } i = 1, \dots, n - 1.$$

Note that this is the same as saying that $H(q, n)$ is the algebra generated by the elements g_i satisfying both the relation given above and the relations on the definition of a braid group, i.e.,

$$g_i g_j = g_j g_i, \text{ for } |i - j| \geq 2$$

$$g_i g_{i+1} g_i = g_{i+1} g_i g_{i+1}, \text{ for } i = 1, \dots, n - 2$$

It is interesting to note that when $q = 1$, we have that $g_i^2 = 1$, which means that, $H(1, q)$ satisfies the relations of the group algebra of the symmetric group S_n , i.e., $H(1, q) = \mathbb{C}[S_n]$, thus, we say that the family of $\mathbb{C}$ -algebras $H(q, n)$, is a *one parameter deformation* of $\mathbb{C}[S_n]$. More than that, it is shown that for q sufficiently next to 1, these algebras are isomorphic to $\mathbb{C}[S_n]$, for such q 's, we have that $H(q, n)$ are semisimple, just like $\mathbb{C}[S_n]$, i.e., they can be written as a direct sum (of simple algebras, i.e., those whose only ideals are the trivial ones) of matrix algebras.

In [12], it is established that a useful basis for a Hecke algebra $H(q, n)$ is given by:

Proposition 5.4.1. *The set*

$\{(g_{i_1} g_{i_1-1} \cdots g_{i_1-k_1})(g_{i_2} g_{i_2-1} \cdots g_{i_2-k_2}) \cdots (g_{i_p} g_{i_p-1} \cdots g_{i_p-k_p}); 1 \leq i_1 < i_2 < \dots < i_p \leq n-1\}$
forms a basis for the $\mathbb{C}$ -module $H(q, n)$, for a generic q .

With this proposition, we are now able to study Markov traces over the Hecke algebras. Note that since $B_n \subset \mathbb{C}[B_n]$ then the map that relates each generator of $\mathbb{C}[B_n]$ to a generator of $H(q, n)$ (i.e., the map $\sigma_i \mapsto g_i$) implies that for each representation of $H(q, n)$, i.e., for each map $H(q, n) \rightarrow \text{GL}(\Lambda)$, there is a representation of B_n through the canonical projection.

5.4.2 Birman-Murakami-Wenzl Algebras

Jun Murakami, Joan Birman and H. Wenzl, independently introduced a family of finite $\mathbb{C}$ -algebras of two parameters, denoted by $C_n(\alpha, l)$, where α and l are complex numbers satisfying $\alpha^4 \neq 1$ and $l^4 \neq 1$. For $j = 1, \dots, n-1$, we call

$$e_j = (\alpha + \alpha^{-1})^{-1}(\sigma_j + \sigma_j^{-1}) - 1 \in \mathbb{C}[B_n].$$

These algebras are also consequences of the study of representations of the braid group B_n . The algebra $C_n(\alpha, l)$, is then defined as the quotient of the group algebra $\mathbb{C}[B_n]$ with the relations:

$$e_j \sigma_j = l^{-1} e_j, \quad e_j \sigma_j^{\pm 1} e_j = l^{\pm 1} e_j.$$

As before, we may see $C_n(\alpha, l)$ as the group algebra $\mathbb{C}[B_n]$, but satisfying not only it's usual relations, but also the relations mentioned above.

Both the algebraic structure, and the representations of $C_n(\alpha, l)$ were studied by Wenzl, who noticed that when α and l are not roots of unity, and il (here $i^2 = \sqrt{-1}$) is not an integer power of $-i\alpha$, the algebra $C_n(\alpha, l)$ is semisimple.

5.5 The HOMFLY Polynomial

Theorem 5.5.1. *For any complex number $z \in \mathbb{C}$, there is a unique Markov trace on the Hecke algebras $H(q, n)$, satisfying:*

$$t(ab) = t(ba), \quad \forall a, b \in H(q, n)$$

$$t(1) = 1$$

$$t(ag_n) = zt(a), \quad \forall a \in H(q, n)$$

Proof. The idea of the proof will be the following:

- We first show by induction that the words on $H(q, n+1)$ can be written as linear combinations of words of the form $xg_n y$ with $x, y \in H(q, n)$ and words in $H(q, n)$.
- Then, we suppose that there is a Markov trace in $H(q, n)$ and show that we can extend it to $H(q, n+1)$.
- Notice that showing that the extension given in the previous item is a Markov trace is equivalent to showing that it satisfies $t(xg_n y g_n) = t(g_n x g_n y)$ for all $x, y \in H(q, n)$.

We shall do this now! Note that from the quadratic relation we may change any negative exponent of some g_j by a sum of elements with non negative exponents, by using the relation $g_j^{-1} = \frac{1}{q}g_j - \frac{q-1}{q}$. We show now that every element in $H(q, n+1)$ may be written as a linear combination of elements of the form $xg_n y$ with $x, y \in H(q, n)$ words in $H(q, n)$. In fact, the result is clearly true for $n = 2$, since the words in $H(q, 1)$ are complex numbers (since g_1 is the identity). Now, suppose the result holds for $H(q, n)$ then taking any arbitrary word in $H(q, n+1)$ in which g_n shows up two times, such a word is given by $g_n x g_n$ with $x \in H(q, n)$, now, we have that $x = y g_{n-1} z$ with $y, z \in H(q, n-1)$ or $x \in H(q, n-1)$. The second case is simple, since $x \in H(q, n-1)$ and $g_n \in H(q, n+1)$,

then they commute, then the quadratic relation simplifies the expression and we get to our result. In the first case, we have:

$$g_n x g_n = g_n (y g_{n-1} z) = y (g_n g_{n-1} g_n) z = y (g_{n-1} g_n g_{n-1}) z = (y g_{n-1}) g_n (g_{n-1} z)$$

as we wanted.

Suppose now that a Markov trace exists in $H(q, n)$ and we shall extend it to $H(q, n+1)$. Consider $t: H(q, n+1) \rightarrow \mathbb{C}$, given by

$$\text{tr}(x g_n y) = z \text{tr}(xy), \quad \forall x, y \in H(q, n)$$

then tr defines a linear functional in $H(q, n+1)$, so it's sufficient to show that tr is a Markov trace. For such, notice that it is enough to show that $\text{tr}(xy) = \text{tr}(yx)$ for an arbitrary choice of $x \in H(q, n+1)$ and $y \in H(q, n) \cup \{g_n\}$, since the set $H(q, n) \cup \{g_n\}$ generates $H(q, n+1)$, it then follows, according to what we have proved previously, that it is enough to show that

$$\text{tr}(x g_n y g_n) = \text{tr}(g_n x g_n y), \quad \forall x, y \in H(q, n).$$

Let us separate the proof into four cases:

Case 1. If $x, y \in H(q, n-1)$, then the result is immediate, since $g_n \in H(q, n+1)$ and thus, commutes with x and y .

Case 2. If $x = a g_{n-1} b$, with $a, b, y \in H(q, n-1)$, note that

$$\begin{aligned} g_n(a g_{n-1} b) g_n y &= a g_n g_{n-1} g_n b y \\ &= a g_{n-1} g_n g_{n-1} b y \end{aligned}$$

Now, applying the linear functional tr (and remembering that q is a complex scalar):

$$\text{tr}(g_n(a g_{n-1} b) g_n y) = z \text{tr}(a g_{n-1}^2 b y) = (z^2(q-1) + zq) \text{tr}(a b y)$$

In the same fashion we have that:

$$(a g_{n-1} b) g_n y g_n = a g_{n-1} b g_n^2 y$$

Now, applying the linear functional tr , we get:

$$\begin{aligned} \text{tr}(a g_{n-1} b g_n^2 y) &= (1-q)z \text{tr}(a g_{n-1} b g_n y) + q \text{tr}(a g_{n-1} b y) \\ &= z^2(q-1) + zq \text{tr}(a b y) \end{aligned}$$

Since $t(x b_n y b_n) = t(b_n x b_n y)$, we have proved the result for this case. Observe the subtlety in the previous calculations, comes from the fact that, since $a, b, y \in H(q, n-1)$ e $g_n \in H(q, n+1)$, we may commute them freely in the computations.

Case 3. If $y = a g_n b$ and $a, b x \in H(q, n-1)$ we have a completely analogous situation to that we proved before.

Case 4. If $x = ag_{n-1}c$ and $y = fg_{n-1}d$ with $a, b, c, d \in H(q, n-1)$. Note that

$$\begin{aligned} xg_nyg_n &= (ag_{n-1}b)g_n(fg_{n-1}d)g_n = (ag_ng_{n-1}g_nbfg_{n-1}d) \\ &= ag_{n-1}g_ng_{n-1}bfg_{n-1}d \end{aligned}$$

Now, applying the linear functional tr , we get:

$$\begin{aligned} t(ag_{n-1}g_ng_{n-1}bfg_{n-1}d) &= z \text{tr}(ag_{n-1}^2bfg_{n-1}d) \\ &= z(q-1) \text{tr}(ag_{n-1}bfg_{n-1}d) + zq \text{tr}(abfg_{n-1}d) \\ &= z(q-1) \text{tr}(ag_{n-1}bfg_{n-1}d) + z^2q \text{tr}(abfd) \end{aligned}$$

Now, we calculations similar to the previous one, but now for $(ag_{n-1}b)g_nf_{n-1}dg_n$, we see that $\text{tr}(xg_nyg_n) = \text{tr}(g_nxg_ny)$ and therefore, the result follows. $\square$

In the previous proof, note that the Markov trace has also something to do with $\text{tr}(ag_n^{-1}) = \bar{z} \text{tr}(a)$, however, since we can take $g_n^{-1} = \frac{1}{q}g_n - \frac{q-1}{q}$ by using the quadratic relation, the proof of this case is analogous to the one we just did by taking $\bar{z} = \frac{z-q+1}{q}$.

Corolary 5.5.1. Taking $\lambda = \frac{1-q+z}{qz}$, consider L a link represented by the closure of a braid $x \in B_n$. Then

$$X_L(q, \lambda) = \left(-\frac{(1-\lambda q)}{\sqrt{\lambda}(1-q)} \right)^{n-1} (\sqrt{\lambda})^{e(x)} t(x)$$

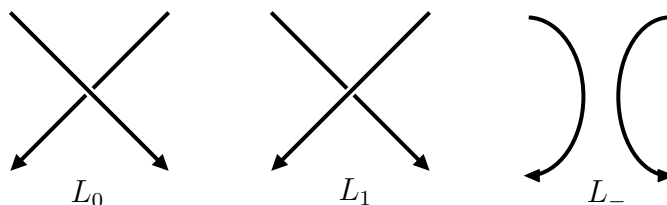
is a link invariant.

The invariant introduced before is called the *HOMFLY polynomial*. It's most common use is through the change of parameter $t = \sqrt{\lambda}q$ and $x = \sqrt{q} - \frac{1}{\sqrt{q}}$, we then denote

$$P_L(t, x) = X_L(q, \lambda).$$

5.6 HOMFLY polynomial - Properties

Let's start by recalling some notation:



Also, remember that the HOMFLY polynomial is given by the expression:

$$X_L(q, \lambda) = \left(-\frac{(1-\lambda q)}{\sqrt{\lambda}(1-q)} \right)^{n-1} (\sqrt{\lambda})^{e(x)} \text{tr}(x)$$

we shall denote $\xi = \left(-\frac{(1-\lambda q)}{\sqrt{\lambda}(1-q)} \right)^{n-1}$ by simplicity, since this term will not be changed with the manipulations we will do. Furthermore, X_L may also be denoted by $P_L(x, t)$ by taking $t = \sqrt{\lambda q}$ e $x = \sqrt{q} - \frac{1}{\sqrt{q}}$.

Theorem 5.6.1. $P_L(t, x)$ is the only Laurent polynomial in the variables t and x satisfying the relation

$$t^{-1}P_{L_0}(t, x) - tP_{L_1}(t, x) = xP_{L_-}(t, x)$$

and assuming the value 1 in $\bigcirc$.

Proof. Using the Markov moves, we can rewrite L_0 , L_1 and L_- as closures of braids in following way L_0 as $\widehat{x\sigma_j^2}$, L_1 as $\widehat{x}$ and L_- as $\widehat{x\sigma_j}$. Now, from the properties of Markov traces, we have:

$$\text{tr}(x\sigma_j^2) = \text{tr}(x[(1-q)\sigma_j + q]) \Rightarrow \text{tr}(x\sigma_j^2) = (q-1)x\text{tr}(x) + q\text{tr}(x).$$

We shall use this notation to prove that P_L satisfies the relations stated above. To do so, we will use the expression X_L previously stated. we have the following:

$$\begin{aligned} X_{L_0}(q, \lambda) &= \xi(\sqrt{\lambda})^{e(x\sigma_j^2)} \text{tr}(x\sigma_j^2) \\ &= \xi(\sqrt{\lambda})^{e(x)+2} [\text{tr}(x\sigma_j)(q-1) + q\text{tr}(x)] \\ &= \lambda\xi(\sqrt{\lambda})^{e(x)} \text{tr}(x\sigma_j)(q-1) + \lambda\xi(\sqrt{\lambda})^{e(x)} \text{tr}(x) \\ &= \xi(\sqrt{\lambda})^{e(x)} \text{tr}(x) (\lambda qz - \lambda z + \lambda q) \end{aligned}$$

Now, since L_1 is the closure $\widehat{x}$, we get

$$X_{L_1}(q, \lambda) = \xi(\sqrt{\lambda})^{e(x)} \text{tr}(x)$$

also, for X_{L_-} we have:

$$\begin{aligned} X_{L_-}(q, \lambda) &= \xi(\sqrt{\lambda})^{e(x\sigma_j)} \text{tr}(x\sigma_j) \\ &= \xi(\sqrt{\lambda})^{e(x)+1} z \text{tr}(x) \end{aligned}$$

Recall now that $t = \sqrt{\lambda}\sqrt{q}$ then $t^{-1} = \frac{1}{\sqrt{\lambda}\sqrt{q}}$, thus we get:

$$\begin{aligned} t^{-1}P_{L_0}(t, x) - tP_{L_1}(t, x) &= \frac{1}{\sqrt{\lambda}\sqrt{q}} \left[\xi(\sqrt{\lambda})^{e(x)} \text{tr}(x) \right] (\lambda qz - \lambda z + \lambda q) - \sqrt{\lambda}\sqrt{q} (\xi(\sqrt{\lambda})^{e(x)} \text{tr}(x)) \\ &= \xi(\sqrt{\lambda})^{e(x)+1} \text{tr}(x) \sqrt{q} \left(z - \frac{z}{q} + 1 - 1 \right) \\ &= \xi(\sqrt{\lambda})^{e(x)+1} \text{tr}(x) \sqrt{q} z \left(1 - \frac{1}{q} \right) \\ &= \xi(\sqrt{\lambda})^{e(x)+1} z \text{tr}(x) \left(\sqrt{q} - \frac{1}{\sqrt{q}} \right) \\ &= \xi(\sqrt{\lambda})^{e(x)+1} z \text{tr}(x) x = X_{L_-}(q, \lambda) x. \end{aligned}$$

Now, we only have to prove uniqueness. Let's prove it by using induction on the number of crossings of L , in the case where $L = \bigcirc$ is trivial, we already have that

$P_{\bigcirc}(x, t) = 1$. suppose now that for a link L with n crossings, the polynomial is unique, i.e, any two polynomials satisfying the relations must be equal, for L with n crossings.

Now, suppose L with $n + 1$ crossings, what the relations above tell us is that the polynomial P_L may be found by using the polynomial of the link L with one crossing changed and the link L with that same crossing smoothed. It is clear that when changing one crossing of a link, we reduce it's number of crossings, as well as smoothing a crossing. This means that the polynomials may be computed using only polynomials of links with a smaller number of crossings, it follows then, by the induction hypothesis that two polynomials of the same link L of $n + 1$ crossings satisfying the above relation, must be equal, hence the HOMFLY polynomial is unique. $\square$

Let's take another look at the expression of the HOMFLY polynomial:

$$X_L(q, \lambda) = \left(-\frac{1 - q\lambda}{\sqrt{\lambda}(1 - q)} \right)^{n-1} (\sqrt{\lambda})^{e(x)} \text{tr}(x)$$

if we try to compute this polynomial in values of q and λ in which the polynomial is not defined, we obtain the *Alexander-Conway polynomial* $\nabla_L(t)$. Thus, such polynomial is given by:

$$\nabla_L(t) = X_L\left(t, \frac{1}{t}\right) = P_L\left(1, \sqrt{t} - \frac{1}{\sqrt{t}}\right).$$

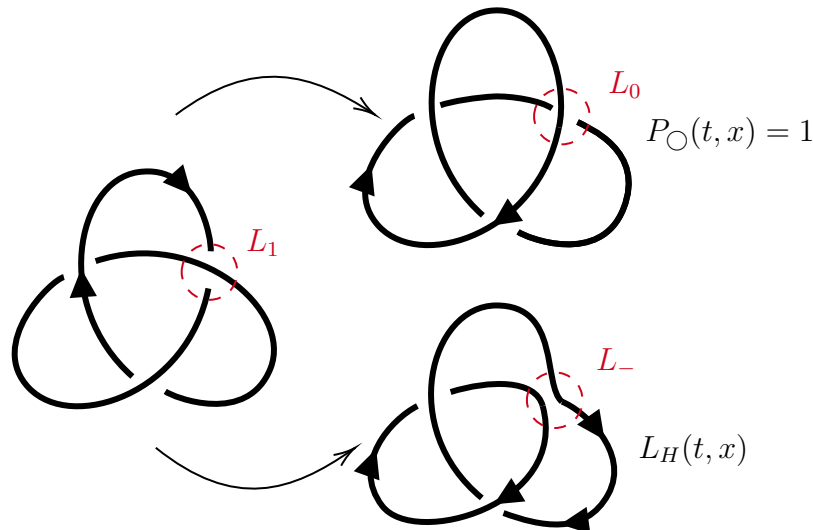
Also, another case where the HOMFLY polynomial cannot be computed is the case $X_L(t, t)$, or $P_L\left(\frac{i}{t}, i\left(\sqrt{t} - \frac{1}{\sqrt{t}}\right)\right)$, where $i \in \mathbb{C}$, the polynomial obtained in such a way is the already known Jones polynomial, since it satisfy the same skein relations, and as we've seen, the Jones polynomial is the only one satisfying such relations, i.e.

$$X_L(t, t) = P_L\left(\frac{i}{t}, i\left(\sqrt{t} - \frac{1}{\sqrt{t}}\right)\right) = V_L(t).$$

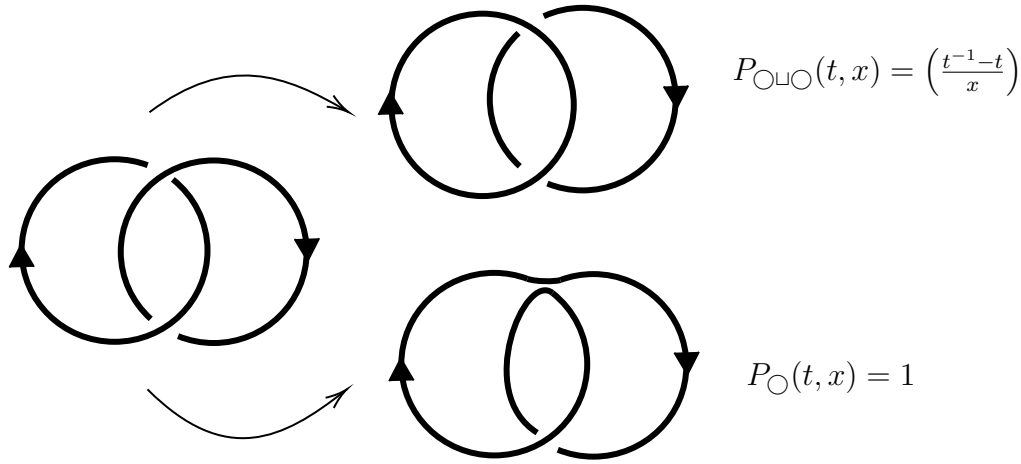
Example. The HOMFLY polynomial of the trefoil (3_1) is given by

$$X_{3_1}(q, \lambda) = \lambda(1 + q^2 - \lambda q^2) \text{ ou } P_{3_1}(t, x) = 2t^2 - t^4 + t^2x^2.$$

Let's call H the diagram of the Hopf link, observe the following diagram:



From this diagram we have that $t^{-1}P_{3_1}(t, x) - tP_{\bigcirc} = xP_H(t, x)$. Thus, to find the HOMFLY polynomial of 3_1 , we first need to find the polynomial of the Hopf link. Again, observe the following diagram:



We have that:

$$t^{-1}P_H(t, x) - tP_{\bigcirc \sqcup \bigcirc}(t, x) = xP_{\bigcirc}(t, x)$$

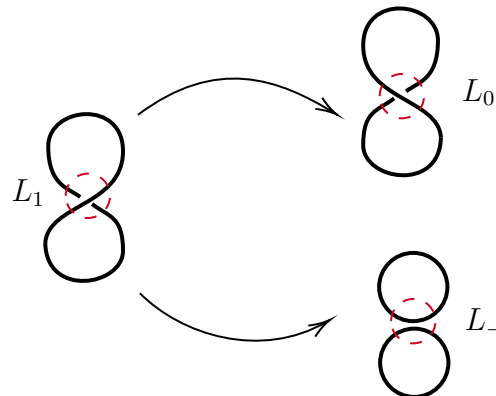
hence, $P_H(x, t) = xt + \left(\frac{t - t^3}{x}\right)$ and putting this back on the skein relation for $P_{3_1}(t, x)$, we get:

$$\begin{aligned} t^{-1}P_{3_1}(t, x) - t &= x \left[xt + \left(\frac{t - t^3}{x}\right) \right] \\ \Rightarrow t^{-1}P_{3_1}(t, x) &= x^2t + 2t = t^3 \\ \Rightarrow P_{3_1}(t, x) &= x^2t^2 + 2t^2 - t^4. \end{aligned}$$

If $\bigsqcup_k \bigcirc$ is the trivial link with k components, then we may compute $P_{\bigsqcup_k \bigcirc}(t, x)$, in two ways:

- First, by noting that $\bigsqcup_k \bigcirc$ is the closure of the identity in B_k ;
- The second, is using the trivial link with $k - 1$ crossings and the skein relations.

We shall do the case $\bigcirc \sqcup \bigcirc$. Consider the figure below, which represents the three diagrams that will be used on the skein relation:



$$\begin{aligned} xP_{L_-}(t, x) &= t^{-1}P_{L_0}(t, x) - tP_{\infty}(t, x) \\ \Rightarrow xP_{\bigcirc \sqcup \bigcirc}(t, x) &= t^{-1} - t \\ \therefore P_{\bigcirc \sqcup \bigcirc}(t, x) &= \frac{t^{-1} - t}{x} \end{aligned}$$
$$P_{\sqcup_k \circ}(t, x) = \left(\frac{t^{-1} - t}{x} \right)^{k-1}.$$

A diagram showing a 2D grid with two rectangles. The top rectangle is labeled x and the bottom rectangle is labeled y . The grid is composed of vertical and horizontal lines. The rectangle x is positioned in the upper left, and rectangle y is positioned in the lower right. Ellipses ($\dots$) are used to indicate that the grid continues beyond the boundaries of the rectangles.

Suppose now, as induction hypothesis, that $\text{tr}(ww') = \text{tr}(w)\text{tr}(w')$ for $w \in B_n$ and $w' \in B_{k-2}$. Thus, taking $w'' \in B_{k-1}$, it follows that $w'' = w'\sigma_{n-1}^\varepsilon$, for some $w' \in B_{n-2}$ and some $\varepsilon \in \mathbb{Z}$. Note that, since even if w'' contains some power of σ_{n-1} that is not at the

right of a word $w' \in B_{n-2}$, when we compute the trace of this word, we can commute the generators in such a way that we obtain the word w'' as written previously. Thus, we get the following:

$$\mathrm{tr}(ww'') = \mathrm{tr}(ww'\sigma_{n-1}^\varepsilon) = z^\varepsilon \mathrm{tr}(ww')$$

now, using the induction hypothesis we obtain $\mathrm{tr}(ww') = \mathrm{tr}(w) \mathrm{tr}(w')$ hence:

$$\mathrm{tr}(ww'') = z^\varepsilon \mathrm{tr}(w') \mathrm{tr}(w) = \mathrm{tr}(w'\sigma_{n-1}) \mathrm{tr}(w) = \mathrm{tr}(w'') \mathrm{tr}(w).$$

Finally, we just have to note that the braid whose closure is $L_1 \# L_2$ is exactly xy , with x and y defined as in the beginning of this proof, hence $e(xy) = e(x) + e(y)$ and, by the observation we did, $\mathrm{tr}(xy) = \mathrm{tr}(x) \mathrm{tr}(y)$, thus, it follows from the expression of $X_L(q, \lambda)$ that:

$$\begin{aligned} X_{L_1 \# L_2}(q, \lambda) &= \left(-\frac{1 - \lambda q}{\sqrt{\lambda}(1 - q)} \right)^{n-k-2} (\sqrt{\lambda})^{e(xy)} \mathrm{tr}(xy) \\ &= \left[\left(-\frac{1 - \lambda q}{\sqrt{\lambda}(1 - q)} \right)^{n-1} (\sqrt{\lambda})^{e(x)} \mathrm{tr}(x) \right] \left[\left(-\frac{1 - \lambda q}{\sqrt{\lambda}(1 - q)} \right)^{k-1} (\sqrt{\lambda})^{e(y)} \mathrm{tr}(y) \right] \\ &= X_{L_1}(q, \lambda) X_{L_2}(q, \lambda) \end{aligned}$$

as we wanted. $\square$

Proposition 5.6.1. *The HOMFLY polynomial is invariant under changing orientation of all the components of a link.*

Proof. Note that changing the orientation of all components of the link does not change the type of crossings. Thus, the skein relation is the same for the new link with the orientation reversed and, therefore, the HOMFLY polynomial is invariant, by using its uniqueness property. $\square$

Proposition 5.6.2. *Let $\bar{L}$ be the mirror image of the link L , then*

$$X_{\bar{L}}(q, \lambda) = X_L\left(\frac{1}{q}, \frac{1}{\lambda}\right) \text{ and } P_{\bar{L}}(t, x) = P_L\left(\frac{1}{t}, -x\right).$$

Proof. If we use the substitutions:

$$\begin{cases} t \mapsto t^{-1} \\ x \mapsto -x \end{cases}$$

we obtain the following skein relation:

$$tP_{L_0}(t^{-1}, -x) - t^{-1}P_{L_1}(t^{-1}, -x) = -xP_{L_-}(t^{-1}, -x)$$

on the other hand, the mirror image of L is obtained by changing each crossing of L . Thus, changing L_0 by L_1 and L_- is kept the same, then from the original skein relation, we obtain:

$$t^{-1}P_{L_1}(t, x) - t^{-1}P_{L_0}(t, x) = xP_{L_-}(t, x) \Rightarrow tP_{L_0}(t, x) - t^{-1}P_{L_1}(t, x) = -xP_{L_-}(t, x)$$

Which means that both polynomials satisfy the same skein relation and, therefore, from the uniqueness property, they must be the same, i.e.,

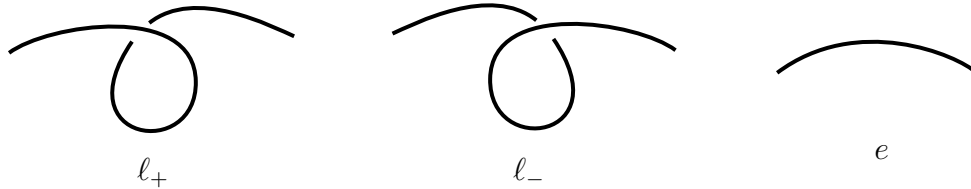
$$P_{\bar{L}}(t, x) = P_L\left(\frac{1}{t}, -x\right).$$

the change of variable that gives us the expression of X_L given P_L guarantees the second equality. $\square$

5.7 The Kauffman Polynomial

We now define a theorem and a corollary that, together, define the Kauffman polynomial. However, to be able to prove them, we are going to need some more tools. The goal of this section will be to introduce not only these tools but also, the Kauffman polynomial, which is also a link invariant.

Theorem 5.7.1. *Let's start by establishing the following notations:*



There is only one function Λ defined on the set of non oriented diagrams taking values in $\mathbb{Z}[a^{\pm 1}, z^{\pm 1}]$ and satisfying the following conditions:

- (i) $\Lambda(\bigcirc) = 1$.
- (ii) Λ is invariant through Reidemeister moves R_2 and R_3 .
- (iii) $\Lambda(\ell_+) = a\Lambda(e)$ e $\Lambda(\ell_-) = a^{-1}\Lambda(e)$.
- (iv) Λ satisfy the following skein relation:

$$\Lambda(L_1) + \Lambda(L_0) = z(\Lambda(L_-) + \Lambda(L_+)).$$

Corolary 5.7.1. *Let $\mathbf{D}$ be an oriented diagram for the link L and $w(\mathbf{D})$ its writhe. Then*

$$K_L(a, z) = a^{w(\mathbf{D})}\Lambda(D) \in \mathbb{Z}[a^{\pm 1}, z^{\pm 1}]$$

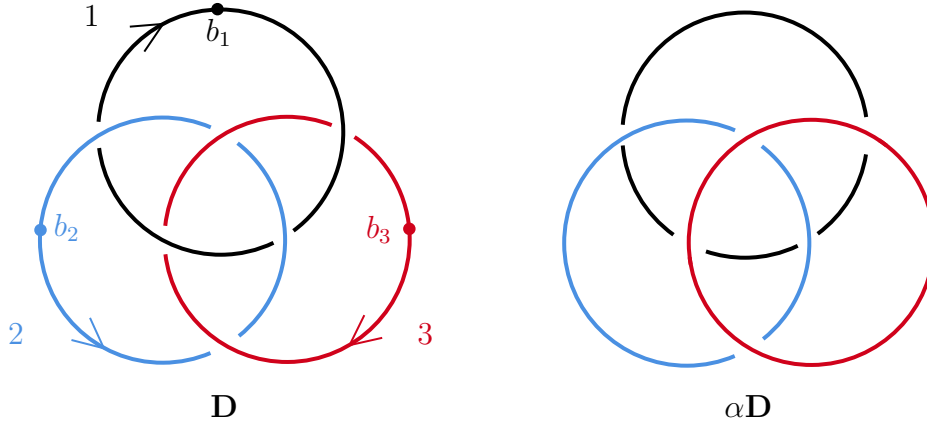
is a link invariant.

In order to prove such theorem we introduce some new concepts. We say a diagram D of a link L is *ordered* if there is a well defined order in the set of components of L . Furthermore, D is *pointed* if for each component a base point is chosen.

We define the *ascending diagram* of an ordered, oriented and pointed diagram $\mathbf{D}$, and denote it by $\alpha\mathbf{D}$, as the diagram formed by $\mathbf{D}$ in the following way: Take the first component of $\mathbf{D}$ and walk along it following the direction given by its orientation and starting at its base point. We change the crossings of this diagram (if necessary) in such a way that the first encounter with each crossing is an undercrossing.

It is important to note that each crossing shall be changed only once even if it is reached again when working on another component of the same link. Thus, it can happen that when we reach a crossing again, it is not an undercrossing, however, we shall not change it again, since this is not the first time we reach it.

The following is the ascending diagram for the Borromean rings link:



Now, we introduce a slight generalization for the ascending diagram. We first define the *untying function*, which, as we shall see, can also be thought of as the height function for some kinds of links.

Definition 5.7.1. The function $h: D \rightarrow \mathbb{R}$ where D is a diagram for an ordered and pointed link L . h is called the *untying function* when it is a two-valued function at the crossings and satisfies the following conditions:

- (i) If the components $c_i < c_j$, then $h(x_i) < h(x_j)$ for any $x_i \in c_i$ and $x_j \in c_j$;
- (ii) In each component c_i of the link L , the function h is monotone strictly increasing from the base point b_i until a certain top point t_i , in any direction.
- (iii) In each crossing, the value of h at the overcrossing is greater than the value of h at the undercrossing.

It is important to note that every ascending diagram has an untying, to see this, we just need to consider, in each component, the function that is increasing from the base point to a certain top point, and such that the top point of a component is immediately smaller than the base point of the next component. Also, note that if a diagram has an untying function, then the image of this diagram is a trivial link with some number of components. Thus, each height achieved by this link's height function (which, in this case, coincides with the untying function) is reached at most two times due to the monotonicity of h . Consequently, the set of segments that connect points of same height is a disjoint union of bounded disks having as boundary the components of L .

With this in mind, we also note that not every link has an untying function, and a very simple counter example is the trefoil knot.

Observation. We change a little bit the notations now, we will use:

$$L_1 = D_+ \quad L_0 = D_- \quad L_- = D_0 \quad \text{e} \quad L_+ = D_\infty$$

Now, let's note the following:

- In Theorem 5.7.1 note that if the diagram D of the link L is non oriented, then there is no correct way to determine if the crossings are of type D_+ or D_- . However, note that interchanging these crossings in the skein relation, we still obtain the same relation, i.e., these crossings are symmetric in the skein relation.
- The same type of symmetry occurs with D_0 and D_∞ .

- A solution for the skein relation presented on the theorem is the ordered 4-tuple $(\Lambda(D_+), \Lambda(D_-), \Lambda(D_0), \Lambda(D_\infty)) = (ax, a^{-1}x, x, \delta x)$, onde $\delta := \frac{a + a^{-1}}{z}$.

In fact, the third item presented above is true due to:

$$\begin{aligned} ax + a^{-1}x &= z \left(x + \left(\frac{a + a^{-1}}{z} \right) x - x \right) \\ \Rightarrow x(a + a^{-1}) &= (a + a^{-1})x \\ \Rightarrow 0 &= 0 \end{aligned}$$

Now, we shall prove using induction on n that the function $\Lambda: \mathcal{D}_n \rightarrow \mathbb{Z}[a^{\pm 1}, z^{\pm 1}]$ is defined, here $\mathcal{D}_n$ is the set of all diagrams with, at most, n crossings. The induction hypothesis is that such function is defined on $\mathcal{D}_{n-1}$ and satisfies:

- (a) The skein relation presented on the theorem is satisfied for every four diagrams related as in the theorem.
- (b) $\Lambda(\ell_+) = a\Lambda(e)$ and $\Lambda(\ell_-) = a^{-1}\Lambda(e)$;
- (c) $\Lambda(D)$ is invariant via Reidemeister R_2 and R_3 that never involves more than $n - 1$ crossings;
- (d) If D is any diagram in $\mathcal{D}_{n-1}$ with $\#D$ components that have untying function, then $\Lambda(D) = a^{\bar{w}(D)}\delta^{\#D-1}$, where $\bar{w}(D)$ denote the sum of all signs of crossings occurring in the same component, such a number is called the *self-writhe*.

The first induction step is $\mathcal{D}_0$, where the only link is the unknot, that has trivial untying function.

Now, to extend the definition of Λ to $\mathcal{D}_n$, we consider D a diagram with n crossings, we choose an orientation, an order and a set of base points for its components. Thus, consider αD the ascending diagram of D and define: $\Lambda(\alpha D) = a^{\bar{w}(\alpha D)}\delta^{\#D-1}$, note that in an ascending diagram $\#\alpha D = \#D$. The value of $\Lambda(D)$ is then computed from αD using the skein relation on Theorem 5.7.1, changing each crossing of αD , until we get D , note that doing it like described, the order of the crossings we choose to change in αD is not important.

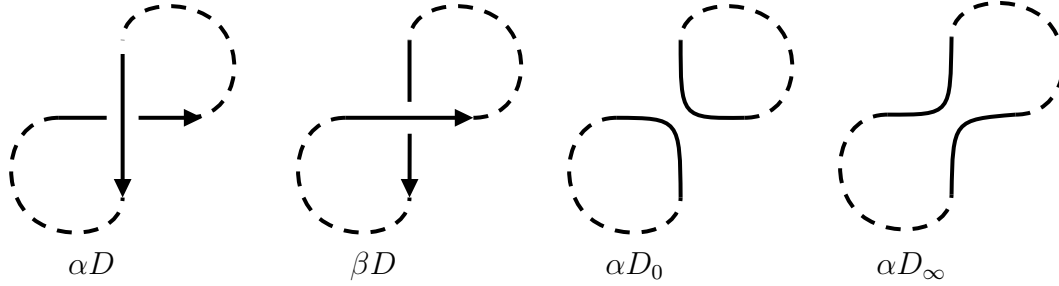
Before we continue, note that in the beginning of our argument, we've done lots of choices regarding D , in fact, we have chosen an order, an orientation and a set of base points. Thus, in order to continue, we shall show that obtaining $\Lambda(D)$, is independent of each of these choices.

Lemma 5.7.1. $\Lambda(D)$ does not depend on the choice of base points.

Proof. We just need to check what happens when we move a base point b of one component, following the direction given by the orientation, until it crosses the first crossing, then we take b' as the new base point, which is the point immediately after this crossing. Note that if the crossing happens between two distinct components, if we call αD the ascending diagram with respect to the base point b and βD the ascending diagram with respect to the base point b' , then both diagrams, in this case, are equal.

However, note that if the crossing happens in the same component, the way we have taken the base point b' and according to the definition of an ascending diagram, it follows

that both diagrams αD and βD are identical, except at the crossing between b and b' , where one of them is an overcrossing, the other one is an undercrossing. Now, let's call αD_0 and αD_∞ the diagrams obtained from αD , by undoing the crossings as in D_0 and in D_∞ (like before).



Now, let's make some observations of what we've obtained:

- Since each created (disjoint) component in the diagram of αD_0 is part of an ascending diagram, then αD_0 is an ascending diagram and, therefore, it has an untying function;
- αD_∞ has an untying function coming from that of αD , modified only on the new arcs that were created by the region where there was a crossing.
- Since αD_0 and αD_∞ have one crossing less, i.e., $\alpha D_0, \alpha D_\infty \in \mathcal{D}_{n-1}$, we may use the induction hypothesis to obtain:

$$\Lambda(\alpha D_0) = a^{\bar{w}(\alpha D_0)} \delta^{\# \alpha D_0 - 1}, \quad \Lambda(\alpha D_\infty) = a^{\bar{w}(\alpha D_\infty)} \delta^{\# \alpha D_\infty - 1}.$$

- Since the ascending diagrams are trivial links, it follows that the linking number between any two of its components is zero and, hence, $\bar{w}(\alpha D_0) = \bar{w}(\alpha D_\infty)$. Furthermore, it is clear that $\# \alpha D_\infty = \# D$ and $\# \alpha D_0 = \# D + 1$.
- We also have that, if ε is the sign of the crossing that differs from αD to βD , then $\bar{w}(\alpha D) = \bar{w}(\alpha D_0) + \varepsilon$ e $\bar{w}(\beta D) = \bar{w}(\alpha D_0) - \varepsilon$.

Having all this information in hand, the next step is to compute βD in two different ways, first, we consider b' as the base point, in this case, it suffices to use the definition to get $\Lambda(\beta D) = a^{\bar{w}(\beta D)} \delta^{\# D - 1}$. Second, we use b as the base point and, since βD is obtained from αD by changing only one crossing, it suffices to use, only once, the skein relation to find the value of $\Lambda(\beta D)$. Now, note that substituting in the skein relation:

$$\Lambda(\alpha D) + \Lambda(\beta D) = z(\Lambda(\alpha D_0) + \Lambda(\alpha D_\infty))$$

the value $\Lambda(\beta D) = a^{\bar{w}(\beta D)} \delta^{\# D - 1}$, we obtain:

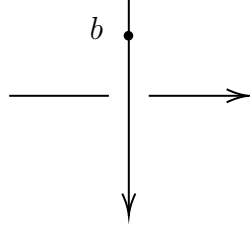
$$a^{\bar{w}(\alpha D_0)} (a \delta^{\# D - 1} + a^{-1} \delta^{\# D - 1}) = z a^{\bar{w}(\alpha D_0)} (\delta^{\# D} + \delta^{\# D - 1})$$

which agrees with the definition given to δ .

Thus, since the value of $\Lambda(\beta D)$ can be obtained from αD or from the definition itself, when we use b' as a base point, it follows that the value of $\Lambda(D)$, is independent of choice of base point we chose. $\square$

Lemma 5.7.2. *The skein relation is valid in $\mathcal{D}_n$.*

Proof. Consider $D = D_+$, let's choose a base point b immediately before the crossing and in such a way that if αD the ascending diagram of D , then the crossing have been changed in αD .



Thus, to obtain $\Lambda(D)$, we've used the skein relation starting in αD , hence, the skein relation $\Lambda(D_+) + \Lambda(D_-) = z(\Lambda(D_0) + \Lambda(D_\infty))$ is precisely the first step in this calculation, using the chosen base point. □

Lemma 5.7.3. *Considering the notations in the beginning, it is true that:*

$$\Lambda(\ell_+) = a\Lambda(e) \text{ e } \Lambda(\ell_-) = a^{-1}\Lambda(e).$$

Proof. Suppose D is a diagram containing a loop of type ℓ_+ and D' a diagram that is identical to D , but that does not contain such loop. Choose a base point b immediately before the crossing on the loop and in such a way that this crossing is an undercrossing so that, when we compute the ascending diagram, the sign of this crossing is not changed. Now, note that the diagram αD is obtained from $\alpha D'$ by adding a loop ℓ_+ , note that this is also valid for each step on the computation of $\Lambda(D)$ starting in αD . Also, depending on the sign of the loop, we have

$$\bar{w}(\alpha D) = \bar{w}(\alpha D') \pm 1$$

in other words, $\bar{w}(\alpha D) = \bar{w}(\alpha D')$ is the sign of the crossing on the loop. Now, from the induction hypothesis, we have that

$$\begin{aligned} \Lambda(\alpha D) &= a^{\bar{w}(\alpha D)} \delta^{\#D-1} \\ &= a^{\bar{w}(\alpha D') \pm 1} \delta^{\#D-1} \\ &= a^{\pm 1} (a^{\bar{w}(\alpha D')} \delta^{\#D-1}) \\ &= a^{\pm 1} \Lambda(\alpha D') \end{aligned}$$

Furthermore, $\Lambda(D)$ is obtained from $\Lambda(\alpha D)$, $\Lambda(D')$ is obtained from $\Lambda(\alpha D')$ and $\Lambda(\alpha D) = a^{\pm 1} \Lambda(\alpha D')$, hence, $\Lambda(D)$ is obtained from $a^{\pm 1} \Lambda(\alpha D')$, which means that since the term $a^{\pm 1}$ is kept in every step of the calculation of $\Lambda(D)$, we have $\Lambda(D) = a^{\pm 1} \Lambda(D')$, with the sign that appears on the exponent depending on the type of loop we considered first. □

Lemma 5.7.4. *Λ is invariant under Reidemeister moves R_2 , R_3 and R_4 , that do not create more than n crossings.*

Proof. Let's consider two diagrams D and D' identical except in a small region, where they differ only by a Reidemeister move R_2 , let's then denote by $D' = D_{R_2}$. we choose a base point on the component that goes under the crossing, thus, these small regions are not changed when we study their associated ascending diagrams in any of the steps of computation of $\Lambda(D)$ and $\Lambda(D')$. In particular, we note that in the process of computation of $\Lambda(D)$ and $\Lambda(D')$, starting in $\Lambda(\alpha D)$ and $\Lambda(\alpha D')$ they differ only by this small region where they are different thanks to the Reidemeister R_2 move.

In this way, note that since $\bar{w}(\alpha D) = \bar{w}(\alpha D')$ as well as $\#\alpha D = \#D = \#D' = \#\alpha D'$, and then, it follows that $\Lambda(\alpha D) = \Lambda(\alpha D')$.

Thus, we note that the next diagram on the process of computation of $\Lambda(D)$ does not depend on any crossing where the R_2 move happens, hence, the value of Λ on the next diagram (just after αD) is the same value as the value of Λ on the diagram just after $\alpha D'$, therefore, these values must be equal. Continuing this way, and from the fact that there is no step where the crossings inside the region where the R_2 move happens are changed, we have that $\Lambda(D) = \Lambda(D')$.

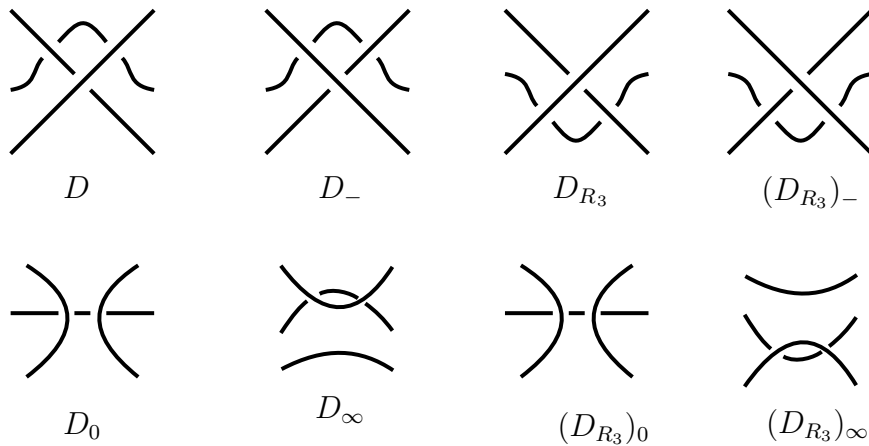
For R_3 , let's consider D and D' which differ only by a small region where the R_3 move happens, denote then $D' = D_{R_3}$. Consider the base points of the components to lie outside this region, and let's call *strand 1* the strand inside the region of the R_3 move that is crossed first, *strand 2* the strand inside the region that is crossed just after strand 1, and similarly, *strand 3* the last strand to be crossed. In other words, we take these nomenclatures to respect the order of the components.

In the ascending diagram, strand 1 is under every other strand, while strand 3 is above all other strands. Thus, from the skein relation, we have:

$$\Lambda(D) + \Lambda(D_-) = z [\Lambda(D_0) + \Lambda(D_\infty)]$$

$$\Lambda(D_{R_3}) + \Lambda((D_{R_3})_-) = z [\Lambda((D_{R_3})_0) + \Lambda((D_{R_3})_\infty)]$$

and note that $\Lambda(D_0) = \Lambda((D_{R_3})_0)$ and $\Lambda(D_\infty) = \Lambda((D_{R_3})_\infty)$, this is seen by observing the following diagrams:



Thus, it follows that, since $D_0 = (D_{R_3})_0$ and D_∞ and $(D_{R_3})_\infty$ are related by a move of type R_2 , considering what was proven previously, we have the relation:

$$\begin{aligned} \Lambda(D) + \Lambda(D_-) &= \Lambda(D_{R_3}) + \Lambda((D_{R_3})_-) \\ \Leftrightarrow \Lambda(D) - \Lambda(D_{R_3}) &= \Lambda(D_-) - \Lambda((D_{R_3})_-) \end{aligned}$$

Now, note that from the above equality, follows that the crossing that is opposite to the strand where the R_3 move is being applied, may be assumed to be of any type (over or undercrossing) since in both cases, the result obtained is the same, we just have to note that:

$$\Lambda(D) + \Lambda((D_{R_3})_-) = \Lambda(D_-) + \Lambda(D_{R_3}).$$

Now, note that choosing an order on the components in such a way that the strand where the R_3 move is being applied, is less than the others, and taking a base point such that the two crossings on this strand with the others are the first to be encountered in the process of computing the ascending diagram, it follows that these two crossings will not be changed. Thus, using these two facts, we note that the R_3 move region can be taken in such a way that is not changed on the computation of the ascending diagram and hence, the proof of this case follows as the previous case.

Note also, we could have considered a case where the R_3 move is applied in a strand that passes over the other strands. In this case, it suffices to choose an order in which the strand that passes under all the others is less than all the other strands, and the base point is to be chosen, again, in a way that we do not change the region of the R_3 move in any of the steps of the computation of the ascending diagram.

For the movement of type R_4 , it suffices to note that, It is a double application of the R_2 move and, therefore, it is immediate that Λ is invariant via R_4 .

□

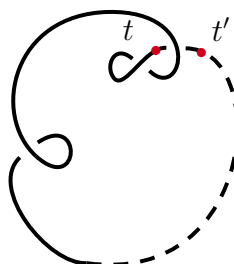
Lemma 5.7.5. *If the diagram $D \in \mathcal{D}_n$ have a height function, then:*

$$\Lambda(D) = a^{\bar{w}(D)} \delta^{\#D-1}$$

Proof. If the top point is the point immediately after the base point in every component, then the diagram is ascending, and from the hypothesis we have that the equality on the lemma is true. We then need to verify what happens when the top point is not the point immediately after the base point, i.e., when there are crossings between the top point and base point.

Consider t to be the top point before a crossing, and take a point t' just after this crossing. No matter if this crossing occurs in the same component or in two different components, if the strand where the points t and t' are, passes over the other, then a small change in the untying function with top point t gives us an untying with top point t' .

Let's consider the case where the strand that passes under the crossing in the same component is the strand that contains the points t and t' . In the following diagram, we represent as a continuous curve a the increasing part of the untying untying function and the dashed part represents the decreasing one.

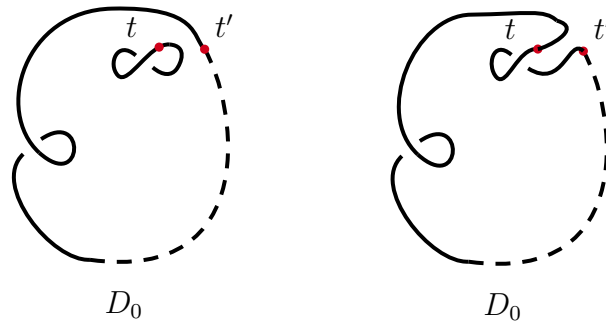


The strand that contains t' is dashed and the other can't be dashed, since it is contained in the part of the diagram that is before t , and therefore, in the ascending part of the diagram.

Consider now the diagram D' that is identical to the diagram D , except by the crossing in discussion, which is changed. From the untying function of D , we make a small change and obtain for D' an untying function. Furthermore, we now have that the crossing just after t is an overcrossing, and we can, then, as previously make a small change on the untying function of D' to obtain t' as the top point of this component, thus, we have that:

$$\Lambda(D') = a^{\bar{w}(D')} \delta^{\#D'-1}.$$

Let's now start to use the skein relation to determine $\Lambda(D)$. Smoothing this crossing in the two possible ways proposed by the relation, we obtain the following diagrams:



Thus, note that $\Lambda(D_0)$ and $\Lambda(D_\infty)$ are known since they have $n - 1$ crossings. More than that, we have that $\bar{w}(D_\infty) = \bar{w}(D_0)$, $\bar{w}(D) = \bar{w}(D_0) + 1$ and $\bar{w}(D') = \bar{w}(D_0) - 1$. Even more, D_0 has one component more than any other diagram involved, which means that, $\#D_0 = \#D + 1$, thus we have:

- $\Lambda(D') = a^{\bar{w}(D_0)-1} \delta^{\#D-1} = a^{\bar{w}(D)-2} \delta^{\#D-1};$
- $\Lambda(D_0) = a^{\bar{w}(D)-1} \delta^{\#D};$
- $\Lambda(D_\infty) = a^{\bar{w}(D)-1} \delta^{\#D-1};$
- $\delta = \left(\frac{a + a^{-1}}{z} \right) - 1 \Rightarrow z = \left(\frac{a + a^{-1}}{\delta + 1} \right).$

And therefore:

$$\begin{aligned} \Lambda(D) + a^{\bar{w}(D)-2} \delta^{\#D-1} &= z \left(a^{\bar{w}(D)-1} \delta^{\#D-1} + a^{\bar{w}(D)-1} \delta^{\#D} \right) \\ \Rightarrow \Lambda(D) + a^{\bar{w}(D)-2} \delta^{\#D-1} &= \left(\frac{a + a^{-1}}{\delta + 1} \right) a^{\bar{w}(D)-1} \delta^{\#D-1} (1 + \delta) \\ \Rightarrow \Lambda(D) + a^{\bar{w}(D)-2} \delta^{\#D-1} &= a^{\bar{w}(D)} \delta^{\#D-1} + a^{\bar{w}(D)-2} \delta^{\#D-1} \\ \therefore \Lambda(D) &= a^{\bar{w}(D)} \delta^{\#D-1} \end{aligned}$$

as we wished.

Lemma 5.7.6. Λ is invariant under orientation of components.

Proof. Let D and αD as usual. After changing the orientation of one of the components of D , we associate a second ascending diagram βD . Of course βD has a height function, the one that comes from the height function of the diagram D with the first choice of orientation of the components. Thus, from the previous lemma, $\Lambda(\beta D) = a^{\bar{w}(\beta D)} \delta^{\#D-1}$. Which means that, we have $\Lambda(\beta D)$ assuming the same value it would assume if it was computed using the second choice of orientation, thus, the computation of $\Lambda(D)$, is independent of the choice of orientation of the components. $\square$

Lemma 5.7.7. *Λ is invariant under the choice of order of the components.*

Proof. Choose D a diagram with a given order and αD its ascending diagram. Choosing a different order for the components and associating to it βD the new ascending diagram with respect to this second choice of ordering, it follows that proving this lemma is the same as proving that, with respect to the first ordering, $\Lambda(\beta D)$ assumes the same value as it would assume with the first choice of ordering, in this case, the result would follow as in the previous lemma.

Let's consider a loop ℓ of the diagram D that is possibly crossed by other strands. If ℓ does not contain any crossing, other than its own, then a Reidemeister move shall be applied in ℓ in such a way that we reduce the height of the link in one unit and, by the induction hypothesis, since βD is ascending, $\Lambda(\beta D) = a^{\bar{w}(\beta D)-1} \delta^{\#D-1}$, which is the value we expected when computing starting at the second choice of ordering. Note also that any component that does not have any crossing in the region that it bounds and such that this region is a disk, may be moved far away from the rest of the link by means of Reidemeister moves of type R_4 .

Now, consider the case where the region bounded by ℓ contains other strands that cross ℓ . Inside the region bounded by ℓ these strands are simple and therefore, the region bounded by one of these strands and the strand that forms ℓ is a 2-gon, which is possibly crossed by more strands of the link. Between all the 2-gons formed in this fashion and also those formed by similar crossings but between two distinct strands contained in the region bounded by ℓ we choose that which is more interior to any other, call it p and q each strand that forms the sides of this 2-gon, call A and B the intersection points between them and call R the region bounded by this 2-gon.

Any other strand that crosses this 2-gon must cross p and q only once and they cross each other at most once in the region R bounded by ℓ . Thus, we can choose a crossing of these two strands in such a way that we have a triangle with a vertex at this crossing and one of its sides in the arc p (or in the arc q) of the 2-gon. from the fact that βD is an ascending diagram, we have that the crossings in this triangle allows us to apply a Reidemeister move of type R_3 , moving the arc p through the crossing chosen previously. The resulting diagram continues to be ascending and we can, again, apply this reasoning.

A finite sequence of applications of this process described on the previous paragraph will allow us to achieve a situation in which the region limited by the 2-gon more interior between all, does not contain any crossing of any other two components. In this case, we can apply a Reidemeister move of type R_2 and reduce the number of crossings of the link to $n - 2$. Hence, using the induction hypothesis we get to the expected result for $\Lambda(\beta D)$. Thus, as in the previous lemmas, the computation of $\Lambda(D)$ is independent of the choice of the ordering of its components, since changing this ordering, the computation of Λ of the new associated ascending diagram (with respect to the new choice of ordering) is the same as the value of Λ of the ascending diagram associated to the first choice of ordering. $\square$

These lemmas, conclude the proof of the main theorem, Theorem 5.7.1, we also have a well defined function Λ satisfying the skein relation given by this theorem and which is invariant via Reidemeister moves R_2 and R_3 . Of course, this does not, yet, implies that Λ is a link invariant, since it is not invariant via the Reidemeister move R_1 . However, we have already encountered a problem similar to this in the past, and the following corollary gives us a way to make use of Λ so that we obtain a link invariant.

Corollary 5.7.2. *Let $\mathbf{D}$ be an oriented diagram for the link L and $w(\mathbf{D})$ its writhe. Then*

$$K_L(a, z) = a^{-w(\mathbf{D})} \Lambda(D) \in \mathbb{Z}[a^{\pm 1}, z^{\pm 1}]$$

is a link invariant.

Proof. We just need to note that

$$w(\ell_+) = w(e) + 1 \text{ e } w(\ell_-) = w(e) - 1$$

while

$$\Lambda(\ell_+) = a\Lambda(e) \text{ e } \Lambda(\ell_-) = a^{-1}\Lambda(e).$$

Thus:

$$K_{\ell_+}(a, z) = a^{-w(\ell_+)} \Lambda(\ell_+) = a^{-w(e)-1} a \Lambda(e) = a^{-w(e)} \Lambda(e) = K_e(a, z)$$

$$K_{\ell_-}(a, z) = a^{-w(\ell_-)} \Lambda(\ell_-) = a^{-w(e)+1} a^{-1} \Lambda(e) = a^{-w(e)} \Lambda(e) = K_e(a, z).$$

□

This invariant is the so called *Kauffman Polynomial*. The Kauffman Polynomial and the Jones Polynomial relate to each other by a change of variable, as indicated in the following proposition:

Proposition 5.7.1. $K_L(-t^{\frac{3}{4}}, t^{-\frac{1}{4}} + t^{\frac{1}{3}}) = V_L(t)$.

Proof. Recall that the Jones polynomial is obtained from the Kauffman bracket, and such bracket satisfies the following properties:

$$\langle D_+ \rangle = A \langle D_0 \rangle + A^{-1} \langle D_\infty \rangle$$

$$\langle D_- \rangle = A^{-1} \langle D_0 \rangle + A \langle D_\infty \rangle$$

i.e., adding these two relations, we obtain:

$$\langle D_+ \rangle + \langle D_- \rangle = (A + A^{-1}) (\langle D_0 \rangle + \langle D_\infty \rangle).$$

Furthermore, other properties that characterize the Kauffman bracket are:

$$\langle \ell_+ \rangle = -A^3 \langle e \rangle$$

but, substituting $(a, z) = (-A^3, A + A^{-1})$, from the skein relations that characterize the polynomial Λ , we have:

$$\Lambda(D_+)(-A^3, A + A^{-1}) + \Lambda(D_-)(-A^3, A + A^{-1}) = (A + A^{-1}) (\Lambda(D_0)(-A^3, A + A^{-1}) + \Lambda(D_\infty)(-A^3, A + A^{-1}))$$

and

$$\Lambda(\ell_+)(-A^3, A + A^{-1}) = -A^3 \Lambda(e)$$

which means that we have,

$$\langle D \rangle = \Lambda(D)(-A^3, A + A^{-1}).$$

Now, from the definitions of V_L and K_L , it immediately follows that:

$$V_L(A) = K_L(-A^3, A + A^{-1})$$

and the change of variables $A = t^{\frac{1}{4}}$ concludes the result. $\square$

To finish, note that if we change all the crossings of a link L , i.e., if we take its mirror image $\bar{L}$, then the skein relation that defined Λ remains the same, however, we have that, since loops of type ℓ_+ become loops of type ℓ_- , then when applying moves of type R_1 in these loops, we obtain that when we had terms a , we now have terms a^{-1} , due to property (iii) of the theorem, i.e., the Kauffman polynomial of the mirror image of a link is given by:

$$K_{\bar{L}}(a, z) = K_L(a^{-1}, z).$$

Furthermore, changing the orientation of every component of a link, the Kauffman polynomial remains the same, since Λ is unchanged by orientation and the writhe of a link with reversed orientation is the same as the one of the original link. $\square$

5.8 The Yang-Baxter Equation

Let's consider V to be the free $\mathbb{K}$ -module of finite type over the commutative ring $\mathbb{K}$, such that $\{e_1, \dots, e_n\}$ is a basis for V . Every endomorphism $R: V \otimes V \rightarrow V \otimes V$ induces an endomorphism

$$R_j: V^{\otimes n} \rightarrow V^{\otimes n}$$

by simply defining R_j to be:

$$R_j = \underbrace{\mathbf{1} \otimes \dots \otimes \mathbf{1}}_{j-1 \text{ times}} \otimes R \otimes \mathbf{1} \dots \otimes \mathbf{1}$$

here, we denote $V^{\otimes n}$ to be the tensor product of n terms V , and $\mathbf{1}$ is the identity endomorphism.

We call the endomorphism R a *Yang-Baxter Operator*, when it satisfies, in $\text{End}(V^{\otimes 3})$ the following identity:

$$R_1 R_2 R_1 = R_2 R_1 R_2$$

which is called the (*quantum*) *Yang-Baxter equation*.

We shall remember now that, if we take a function $f: V^{\otimes n} \rightarrow V^{\otimes n}$, then we can write this function in terms of the generators $\{e_1 \otimes \dots \otimes e_1, \dots, e_n \otimes \dots \otimes e_n\}$ of $V^{\otimes n}$ as:

$$f(e_{i_1} \otimes \dots \otimes e_{i_n}) = \sum_{j=1}^n f_{i_1, \dots, i_n}^{j_1, \dots, j_n} e_{j_1} \otimes \dots \otimes e_{j_n}.$$

Now, let's consider $\mu: V \rightarrow V$ an operator represented by the diagonal matrix $\text{diag}(\mu_1, \dots, \mu_m)$ such that:

$$R \circ (\mu \otimes \mu) = (\mu \otimes \mu) \circ R, \sum_j R_{ij}^{kj} \mu_j = \alpha \beta \delta_i^k, \sum_j (R^{-1})_{ij}^{kj} \mu_j = \alpha^{-1} \beta \delta_i^k$$

where each R_{ij}^{pq} is given by $R(e_i \otimes e_j) = \sum_j R_{ij}^{pq}(e_p \otimes e_q)$, δ_i^k is the Kronecker delta and α and β are fixed unities of $\mathbb{K}$ (recall that a unit is an element of $\mathbb{K}$ which admits an inverse). If such a μ exists, we call it a *enhanced Yang-Baxter operator* (or EYB-operator for short), and denote (R, μ, α, β) when we want to be specific on which units we are using to define μ .

Now, let $d = \dim(V)$, we have that for each $f \in \text{End}(V^{\otimes n})$ we can define a map $\text{tr}_n(f) \in \text{End}(V^{\otimes(n-1)})$ called the *trace operator* defined in the following fashion: Taking $\{v_1, \dots, v_d\}$ a basis for V , then for any $i_1, \dots, i_{n-1} \in \{1, 2, \dots, d\}$:

$$\text{tr}_n(f)(v_{i_1} \otimes \dots \otimes v_{i_{n-1}}) = \sum_{i \leq j_1, \dots, j_{n-1}, j \leq d} f_{i_1, \dots, i_{n-1}, j}^{j_1, \dots, j_{n-1}, j} v_{j_1} \otimes \dots \otimes v_{j_{n-1}}$$

thus, since $\text{tr}(f) = \sum_{1 \leq j_1, \dots, j_n \leq d} f_{i_1, \dots, i_n}^{j_1, \dots, j_n}$ is the trace of f , we must have that $\text{tr}(\text{tr}_n(f)) = \text{tr}(f)$.

To understand what this trace operator does to a function, let's reduce to the case where V is a vector space and g and h are endomorphisms, and let's study $\text{tr}_2(g \otimes h)$. Recall that if we see g and h as matrices:

$$[g] = \begin{pmatrix} g_{11} & \cdots & g_{1m} \\ \vdots & \ddots & \vdots \\ g_{m1} & \cdots & g_{mm} \end{pmatrix}, [h] = \begin{pmatrix} h_{11} & \cdots & h_{1m} \\ \vdots & \ddots & \vdots \\ h_{m1} & \cdots & h_{mm} \end{pmatrix}$$

then, we have:

$$g \otimes h = \begin{pmatrix} g_{11}[h] & \cdots & g_{1m}[h] \\ \vdots & \ddots & \vdots \\ g_{m1}[h] & \cdots & g_{mm}[h] \end{pmatrix}$$

then, if we call $f = g \otimes h$ and $f_{i,j}^{k,l} = g_{ij}h_{kl}$, we have that:

$$[\text{tr}_2(g \otimes h)] = \sum_{i=1}^m g_{ii}[h] = \begin{pmatrix} \sum_i g_{ii}h_{11} & \cdots & \sum_i g_{ii}h_{1m} \\ \vdots & \ddots & \vdots \\ \sum_i g_{ii}h_{m1} & \cdots & \sum_i g_{ii}h_{mm} \end{pmatrix} = \begin{pmatrix} f_{i,1}^{i,1} & \cdots & f_{i,1}^{i,m} \\ \vdots & \ddots & \vdots \\ f_{i,m}^{i,1} & \cdots & f_{i,m}^{i,m} \end{pmatrix}$$

since $\text{tr}_2(g \otimes h)(e_j) = \sum_i f_{i,j}^{i,k} e_k$, which shows us exactly what we wanted to see, but in a particular case, the general one is simply a generalization of notation, hence $\text{tr}(\text{tr}_n(f)) = \text{tr}(f)$.

Furthermore, we still have some additional important properties that the trace operator satisfies. Consider $f \in \text{End}(E^{\otimes(n+1)})$, $g \in \text{End}(V^{\otimes n})$ and $h \in \text{End}(V \otimes V)$, then:

- (i) $\text{tr}_{n+1}(f \circ (g \otimes \text{Id}_V)) = \text{tr}_{n+1}(f) \circ g$;
- (ii) $\text{tr}_{n+1}((g \otimes \text{Id}_V) \circ f) = g \circ \text{tr}_{n+1}(f)$;
- (iii) $\text{tr}_{n+1}(\text{Id}_V^{\otimes(n-1)} \otimes h) = \text{Id}_V^{\otimes(n-1)} \otimes \text{tr}_2(h)$.

A proof of these properties can be found in [23] (note that in this paper Turaev uses $Sp_n = \text{tr}_n$).

By considering $\rho: B_n \rightarrow \text{Aut}(V^{\otimes n})$ a representation of the braid group given in a natural way as $\rho(\sigma_i) = R_i$, we may define a nice link invariant given by the following proposition:

Proposition 5.8.1. *Let (R, μ, α, β) an enhanced Yang-Baxter operator, then the function that maps to each closed braid $\hat{x} \in B_n$, the value:*

$$F(\hat{x}) = \alpha^{-e(x)} \beta^{-n} \text{tr}(\rho(x) \circ \mu^{\otimes n})$$

is a link invariant.

Proof. From the definition of EYB-operator we must have that $\mu^{\otimes n}$ commutes with $\rho(\eta)$ for any $\eta \in B_n$. Thus, for every $\eta, \xi \in B_n$, we have:

$$\text{tr}(\rho(\eta^{-1}\xi\eta) \circ \mu^{\otimes n}) = \text{tr}(\rho(\eta^{-1}) \circ \rho(\xi) \circ \rho(\eta) \circ \mu^{\otimes n}) = \text{tr}(\rho(\eta^{-1}) \circ \rho(\xi) \circ \mu^{\otimes n} \circ \rho(\eta)) \stackrel{(*)}{=} \text{tr}(\rho(\xi) \circ \mu^{\otimes n})$$

where $(*)$ is simply the property of trace $\text{tr}(S \circ T \circ R) = \text{tr}(S \circ R \circ T)$.

Moreover $e(\eta^{-1}\xi\eta) = -e(\eta) + e(\xi) + e(\eta) = e(\xi)$, hence:

$$F(\eta^{-1}\xi\eta) = \alpha^{-e(\eta\xi\eta)} \beta^{-n} \text{tr}(\rho(\eta\xi\eta) \circ \mu^{\otimes n}) = \alpha^{-e(\xi)} \beta^{-n} \text{tr}(\rho(\xi) \circ \mu^{\otimes n}) = F(\xi).$$

Now, we just need to show that $F(\xi\sigma_n^{\pm 1}) = F(\xi)$ for every $\xi \in B_n$ and $\sigma_n \in B_{n+1}$ and the result will follow from the Markov's Theorem. We shall do this now, note that

$$\text{tr}(\rho(\xi\sigma_n) \circ \mu^{\otimes n}) = \text{tr}\left[(\rho(\xi) \otimes \text{Id}_V) \circ R_n \circ (\text{Id}_V^{\otimes(n-1)} \otimes \mu \otimes \mu) \circ (\mu^{\otimes(n-1)} \otimes \text{Id}_V \otimes \text{Id}_V)\right]$$

in fact, we may write $\mu^{\otimes(n+1)}$ as:

$$\mu^{\otimes(n+1)} = \mu^{\otimes(n-1)} \otimes \mu \otimes \mu = (\text{Id}_V^{\otimes(n-1)} \circ \mu^{\otimes(n-1)}) \otimes (\mu \otimes \mu \circ \text{Id}_V \otimes \text{Id}_V) = (\text{Id}_V^{\otimes(n-1)} \otimes \mu \otimes \mu) \circ (\mu^{\otimes(n-1)} \otimes \text{Id}_V^{\otimes 2}).$$

therefore we have:

$$\text{tr}(\rho(\xi\sigma_n) \circ \mu^{\otimes(n+1)}) = \text{tr}\left\{\text{tr}_{n+1}\left[(\rho(\xi) \otimes \text{Id}_V) \circ R \circ (\text{Id}_V^{\otimes(n-1)} \otimes \mu \otimes \mu) \circ (\mu^{\otimes(n-1)} \otimes \text{Id}_V^{\otimes 2})\right]\right\}.$$

And using the properties of trace operator stated previously, we get:

$$\begin{aligned} & \text{tr}_{n+1}\left[(\rho(\xi) \otimes \text{Id}_V) \circ R \circ (\text{Id}_V^{\otimes(n-1)} \otimes \mu \otimes \mu) \circ (\mu^{\otimes(n-1)} \otimes \text{Id}_V^{\otimes 2})\right] \\ &= \text{tr}_{n+1}\left[(\rho(\xi) \otimes \text{Id}_V) \circ (\text{Id}_V^{\otimes(n-1)} \otimes R) \circ (\mu \otimes \mu) \circ (\mu^{\otimes(n-1)} \otimes \text{Id}_V^{\otimes 2})\right] \\ &\stackrel{(i)}{=} \rho(\xi) \circ \text{tr}_{n+1}\left[(\text{Id}_V^{\otimes(n-1)} \otimes R) \circ (\mu \otimes \mu) \circ (\mu^{\otimes(n-1)} \otimes \text{Id}_V^{\otimes 2})\right] \\ &\stackrel{(iii)}{=} \rho(\xi) \circ \left[\text{Id}_V^{\otimes(n-1)} \otimes \text{tr}_2(R \circ (\mu \otimes \mu) \circ (\mu^{\otimes(n-1)} \otimes \text{Id}_V \otimes \text{Id}_V))\right] \\ &\stackrel{(i)}{=} \rho(\xi) \circ \left[\text{Id}_V^{\otimes(n-1)} \otimes \text{tr}_2(R \circ (\mu \otimes \mu))\right] \circ \mu^{\otimes(n-1)} \otimes \text{Id}_V \end{aligned}$$

But, from the definition of EYB-operator, the expression above becomes simply $\alpha\beta(\rho(\xi) \otimes \mu^{\otimes n})$, since $\text{tr}_2(R \circ (\mu \otimes \mu)) = \alpha\beta$, as $\sum_j R_{i,j}^{k,j} \mu_j = \alpha\beta\delta_i^k$, hence we have the following:

$$\text{tr}(\rho(\xi\sigma_n) \circ \mu^{\otimes(n+1)}) = \text{tr}(\alpha\beta\rho(\xi) \otimes \mu^{\otimes n}) = \alpha\beta \text{tr}(\rho(\xi) \otimes \mu^{\otimes n})$$

thus:

$$\begin{aligned} F(\xi\sigma_n) &= \alpha^{-e(\xi\sigma_n)} \beta^{-(n+1)} \text{tr}(\rho(\xi\sigma_n) \otimes \mu^{\otimes(n+1)}) \\ &= \alpha^{-1} \alpha^{-e(\xi)} \beta^{-1} \beta^{-n} \alpha\beta \text{tr}(\rho(\xi) \otimes \mu^{\otimes n}) \\ &= \alpha^{-e(\xi)} \beta^{-n} \text{tr}(\rho(\xi) \otimes \mu^{\otimes n}) = F(\xi). \end{aligned}$$

The other case, $F(\xi\sigma_n^{-1})$ is completely analogous to the one above. Hence, we conclude the proof as the result follows from Markov's Theorem. $\square$

The whole construction and proof above is a demonstration of how important and powerful the Markov's Theorem can be, since we have created a new invariant for links, that, at first glance, looks complicated, but with the help of Markov's Theorem, it turns out that the proof of invariance is actually not surprising.

The only thing that we are missing is to show that

$$\text{tr}_2(R \circ (\mu \otimes \mu)) = \alpha\beta\mu \Rightarrow \sum_{i=1}^m R_{i,j}^{i,l} \mu_i = \alpha\beta\delta_j^l, \forall l \in \{1, \dots, m\}$$

since we used this fact in the proof of the previous proposition.

Recall that if $R = S \otimes T$, then $R \circ (\mu \otimes \mu) = (S \otimes T) \circ (\mu \otimes \mu) = (S \circ \mu) \otimes (T \circ \mu)$. Also, we have:

$$[S \circ \mu] = [S][\mu] = \begin{pmatrix} s_{11} & \cdots & s_{1n} \\ \vdots & \ddots & \vdots \\ s_{m1} & \cdots & s_{mm} \end{pmatrix} \cdot \begin{pmatrix} \mu_1 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & \mu_m \end{pmatrix} = \begin{pmatrix} s_{11}\mu_1 & \cdots & s_{1n}\mu_m \\ \vdots & \ddots & \vdots \\ s_{m1}\mu_1 & \cdots & s_{mm}\mu_m \end{pmatrix}$$

as well as

$$[T \circ \mu] = \begin{pmatrix} s_{11}\mu_1 & \cdots & s_{1n}\mu_m \\ \vdots & \ddots & \vdots \\ s_{m1}\mu_1 & \cdots & s_{mm}\mu_m \end{pmatrix}.$$

Now, we know that:

$$[S \circ \mu] \otimes [T \circ \mu] = \begin{pmatrix} s_{11}\mu_1[T \circ \mu] & \cdots & s_{1n}\mu_m[T \circ \mu] \\ \vdots & \ddots & \vdots \\ s_{m1}\mu_1[T \circ \mu] & \cdots & s_{mm}\mu_m[T \circ \mu] \end{pmatrix}$$

and from the study of the trace operator we did earlier, we get $\text{tr}_2[(S \circ \mu) \otimes (T \circ \mu)] = \sum_{i=1}^m s_{ii}\mu_i[T \circ \mu] \in \text{End}(V)$. Hence, it follows that:

$$[\text{tr}_2[R \circ (\mu \otimes \mu)]] = \begin{pmatrix} \sum_{i=1}^m s_{ii}t_{11}\mu_i\mu_m & \cdots & \sum_{i=1}^m s_{ii}t_{1m}\mu_i\mu_m \\ \vdots & \ddots & \vdots \\ \sum_{i=1}^m s_{ii}t_{m1}\mu_i\mu_m & \cdots & \sum_{i=1}^m s_{ii}t_{mm}\mu_i\mu_m \end{pmatrix} = \begin{pmatrix} \alpha\beta\mu_1 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & \alpha\beta\mu_m \end{pmatrix}$$

and denoting by $R_{i,j}^{k,l} = \delta_i^j t_{kl}$ we have that comparing the two matrices above, we obtain the desired equality:

$$\begin{aligned} \sum_{i=1}^m R_{i,j}^{i,l} \mu_i \mu_l &= \alpha\beta\mu_i \delta_j^l, \forall l \in \{1, \dots, m\} \\ \Leftrightarrow \sum_{i=1}^m R_{i,j}^{i,l} \mu_i &= \alpha\beta\delta_j^l, \forall \{1, \dots, m\}. \end{aligned}$$

5.9 The G_2 Invariant

We now present, in the form of a theorem an invariant of graphs in which triple points are allowed, so far we have been working with two dimensional graphs that are projections of links and carry withing themselves information about over crossings and under crossings, now, we generalize these graphs allowing them to have triple points, this will give us some skein relations similar to those studied on the Kauffman bracket.

The planar graphs we study here will be understood as graphs embedded in S^2 that may have edges that go on a circle around S^2 and have no vertices and also vertices with three edges coming out of it. We may also generalize this even further and allowing more edges to come out of these vertices, but for now we may stick to the three edges case. These graphs will be presented as planar pictures but we may see S^2 as the plane with the point at infinity identified, in a similar fashion we do when thinking about S^3 .

We understand a planar graph, of such kind, to have a boundary, if it is the intersection of a disk D with a planar graph which intersects the boundary of D transversely in a finite number of points, called *endpoints*. The boundary of a graph with boundary is then defined as the boundary of D together with the distinguished endpoints and an inward pointing normal vector, hence this induces us to define the reversed boundary in the following way: let b be the boundary of a graph with boundary then $-b$ is the exact same object but with the normal vector now pointing outwards. The projection of such graphs will be treated the same way as we treated the good projections of links previously. And we will refer to such graphs as *tangled trivalent graphs*, although Kupperberg calls them *freeways* in its paper "*The Quantum G_2 Invariant*".

The importance of this section is to study the construction of such skein relation, and the subsequent construction of an invariant of graphs. This may help us to cook up invariants of 3-manifolds, but this is a topic for the future and is not part of this work.

Lemma 5.9.1. *A crossingless planar graph has at least one simply-connected face with five or fewer sides, which does not share an edge with itself.*

Proof. First, note that if we consider any face with n sides that lies in S^2 , it's Euler characteristic is given by the formula $\chi = V - E + F$, where V is it's number of vertices, E it's number of edges and $F = 1$, which is it's number of faces (which equals 1 since we are already taking only one face). Now, we also note that, since this face has n sides, then $V = E = n$, therefore dividing the number of vertices by 3 and the number of edges by 2, the Euler characteristic of this face can be computed by $\chi = 1 - \frac{n}{6}$. Also, if a face is not simply connected, its Euler characteristic is either 0 or negative.

If we consider a graph with no face that shares an edge with itself, then taking any circle in S^2 will partition S^2 into two disks, and these disks will be partitioned into polygons, since the Euler characteristic of a disk is 1, then any one of the disks we take will have at least one simply connected face, also this face must have at maximum 5 sides since a face with 6 or more sides will contribute with an Euler characteristic less than one, following the previous discussion.

Now, if we suppose a graph with a face that shares an edge with itself, then taking the midpoint of that edge we may create a circle from this point to itself, dividing S^2 into two disks. One of them may have another face that shares an edge with itself, which allows us to repeat the construction. Repeating this process if necessary we can take the innermost circle, in the sense that this circle divides S^2 into two disks and one of them contains only faces that do not share an edge with itself, therefore the same argument can

be used to show that, in this disk, there exists at least one simply connected face with less than 6 sides. $\square$

Now, we construct an invariant for tangled trivalent graphs. Let's first notice that such a construction can be applied to recover the Kauffman bracket, but since we have already done this we refer again to Kuperberg's work to see how the same construction applies to the case of link projections (instead of trivalent graphs projections).

Theorem 5.9.1. *There exists a unique regular isotopy invariant of tangled trivalent graphs called $|\cdot\rangle$ given by the following recurrent relations:*

$$\begin{aligned}
|\emptyset\rangle &= 1 \\
|\bigcirc\rangle &= a |\text{---}\rangle \\
|-\bigcirc\rangle &= 0 \\
|-\bigcirc-\rangle &= b \\
|\text{---}\rangle &= c |\text{---}\rangle \\
|\square\rangle &= d_1 (|\text{---}\rangle + |\text{---}\rangle) + d_2 (|\bigcirc\rangle + |\bigcirc\rangle) \\
|\text{---}\rangle &= e_1 (|\text{---}\rangle + |\text{---}\rangle + |\text{---}\rangle + |\text{---}\rangle) + |\text{---}\rangle \\
&\quad + e_2 (|\text{---}\rangle + |\text{---}\rangle + |\text{---}\rangle + |\text{---}\rangle + |\text{---}\rangle) \\
|\times\rangle &= f_1 |\text{---}\rangle + g_1 |\text{---}\rangle + f_2 |\bigcirc\rangle + g_2 |\bigcirc\rangle
\end{aligned}$$

with the coefficients given by:

$$a = q^5 + q^4 + q + 1 + q^{-1} + q^{-4} + q^{-5}$$

$$b = -q^3 - q^2 - q - q^{-1} - q^{-2} - q^{-3}$$

$$c = q^2 + 1 + q^{-2}$$

$$d_1 = -q - q^{-1}$$

$$d_2 = q + 1 + q^{-1}$$

$$e_1 = 1, e_2 = -1$$

$$f_1 = \frac{1}{1+q}, g_1 = \frac{1}{1+q}$$

$$f_2 = \frac{q}{1+q^{-1}}, g_2 = \frac{q^{-1}}{1+q}.$$

Proof. First we prove the existence. Consider G a tangled graph and $I(G)$ some $\mathbb{K}$ -valued invariant of G , where $\mathbb{K}$ is a field. Now, suppose B is the boundary of G , we can extend $I(G)$ to a vector valued invariant $\vec{I}(G)$ which assumes values in some $\mathbb{K}$ -vector space and depends only on ∂B .

Let $F(b)$ be the vector space of all $\mathbb{K}$ -valued functions defined on the set of all links with boundary $-b$. Then, if $\partial B = b$ we may simply take $\vec{I}(B) \in F(b)$ as being $I(B' \cup B)$, where B' has boundary $\partial B' = -b$, here, we must remember that since $\partial B = b$ and $\partial B' = -b$, then $B' \cup B$ has no boundary, therefore, I is well defined on $B' \cup B$.

Note that $\vec{I}(B) \in F(b)$, but $F(b)$ may be a very large vector space, however, defining the space $\text{Span}(b)$ as the span of $\vec{I}(B)$ for all B with boundary b , then if $\dim(\text{Span}(b))$ is small we may describe the invariant I just in terms of recursive relations on $\vec{I}(B)$ (hence the name recurrent invariant) for only a small number of choices of B , therefore defining I in just a small portion of $F(b)$.

The uniqueness of this invariant follows from the previous lemma, since taking an appropriate choice of parameters to the recurrence relations, the bracket is completely determined by the parameters.

We now proceed by assuming that all acyclic crossingless diagrams with 5 or fewer endpoints are linearly independent, and also, if we take a diagram consisting of a face with k endpoints ($k < 6$), then the bracket of this diagram can be written as a linear combination of those linearly independent diagrams with coefficients given as in the theorem.

It is important to note that on the expression for the square, the first two coefficients are the same due to the symmetry of the shapes, meaning that if we look at the shape vertically and apply the reduction, and then looking at it horizontally and applying the reduction, we get two linear combinations that must be identical, therefore getting the same coefficients as stated. This will also happen for the pentagon, implying that the first five coefficients are all the same, as well as the five last coefficients.

Observe that we have assumed the relations and the coefficients, to complete the proof we have to verify the values of each coefficient, including those of the diagrams with crossings. We begin by computing a, b, c, d_1, d_2, e_1 and e_2 . To do this we will take diagrams that have two adjacent polygons (among triangles, squares, pentagons and circles) and reducing these diagrams in two different manners: suppose for example we have a pentagon adjacent to a square, we first start reducing the pentagon to get a linear combination of the vectors of the basis. Then we start reducing the square, to get another linear combination of the vectors of the basis. Now, since we are assuming that the brackets are invariant under these reductions, the linear combinations must be equal, hence giving us a set of equations.

We start by doing the case exemplified above a pentagon adjacent to a square. Reducing the square first, we get:

$$\begin{aligned}
|\text{pentagon-square}\rangle &= d_1 \left(|\text{pentagon}\rangle + |\text{square}\rangle \right) + d_2 \left(|\text{pentagon}\rangle + |\text{square}\rangle \right) \\
&= d_1 \left(|\text{pentagon}\rangle + d_1 \left(|\text{pentagon}\rangle + |\text{square}\rangle \right) + d_2 \left(|\text{pentagon}\rangle + |\text{square}\rangle \right) \right) + d_2 \left(c |\text{pentagon}\rangle + |\text{square}\rangle \right) \\
&= (d_1 e_2 + d_1 d_2) \left(|\text{pentagon}\rangle + |\text{square}\rangle \right) + (d_1 e_2 + d_2 c) |\text{pentagon}\rangle + (d e_1 + d_2) |\text{square}\rangle \\
&\quad + (d_1^2 + d_1 e_1) \left(|\text{pentagon}\rangle + |\text{square}\rangle \right) + d_1 e_2 \left(|\text{pentagon}\rangle + |\text{square}\rangle \right) + d_1 e_1 \left(|\text{pentagon}\rangle + |\text{square}\rangle \right)
\end{aligned}$$

Now, by reducing first the pentagon, we obtain:

$$\begin{aligned}
|\text{pentagon}\rangle &= e_1 \left(|\text{pentagon}_1\rangle + |\text{pentagon}_2\rangle + |\text{pentagon}_3\rangle + |\text{pentagon}_4\rangle + |\text{pentagon}_5\rangle \right) \\
&+ e_2 \left(|\text{pentagon}_6\rangle + |\text{pentagon}_7\rangle + |\text{pentagon}_8\rangle + |\text{pentagon}_9\rangle + |\text{pentagon}_{10}\rangle \right) + c |\text{pentagon}_{11}\rangle \\
&+ d_1 \left(|\text{pentagon}_{12}\rangle + |\text{pentagon}_{13}\rangle \right) + d_2 \left(|\text{pentagon}_{14}\rangle + |\text{pentagon}_{15}\rangle \right) + d_1 \left(|\text{pentagon}_{16}\rangle + |\text{pentagon}_{17}\rangle \right) + d_2 \left(|\text{pentagon}_{18}\rangle + |\text{pentagon}_{19}\rangle \right) \\
&+ c |\text{pentagon}_{20}\rangle + e_2 \left(b |\text{pentagon}_{21}\rangle + |\text{pentagon}_{22}\rangle + c \left(|\text{pentagon}_{23}\rangle + |\text{pentagon}_{24}\rangle \right) + |\text{pentagon}_{25}\rangle \right) \\
&= (e_1 e_2 + e_2 c) \left(|\text{pentagon}_{26}\rangle + |\text{pentagon}_{27}\rangle \right) + (e_1 e_2 + e_2 b + 2e_1 d_2) |\text{pentagon}_{28}\rangle + e_1^2 |\text{pentagon}_{29}\rangle + \\
&+ (e_1^2 + e_1 d_1 + e_1 c) \left(|\text{pentagon}_{30}\rangle + |\text{pentagon}_{31}\rangle \right) + (e_1 e_2 + e_1 d_2) \left(|\text{pentagon}_{32}\rangle + |\text{pentagon}_{33}\rangle \right) \\
&+ (e_1^2 + e_1 d_1 e_2) \left(|\text{pentagon}_{34}\rangle + |\text{pentagon}_{35}\rangle \right)
\end{aligned}$$

Finally, comparing the two obtained linear combinations, we obtain the following set of equations, which we will not solve yet:

$$\begin{cases}
d_1 e_2 + d_1 d_2 = e_1 e_2 + e_2 c \\
d_1 e_2 + d_2 c = e_1 e_2 + e_2 b + 2e_1 d_2 \\
d_1 e_1 + d_2 = e_1^2 \\
d_1^2 + d_1 e_1 = e_1^2 + e_1 d_1 + e_1 c \\
d_1 e_2 = e_1 e_2 + e_1 d_2 \\
d_1 e_1 = e_1^2 + e_1 d_1 + e_2^2
\end{cases}$$

Next case we must do is the pentagon adjacent to the triangle. Let's first reduce the triangle:

$$\begin{aligned}
|\text{triangle}\rangle &= c |\text{triangle}_1\rangle = c \left[d_1 \left(|\text{triangle}_2\rangle + |\text{triangle}_3\rangle \right) + d_2 \left(|\text{triangle}_4\rangle + |\text{triangle}_5\rangle \right) \right] \\
&= c d_1 \left(|\text{triangle}_2\rangle + |\text{triangle}_3\rangle \right) + c d_2 \left(|\text{triangle}_4\rangle + |\text{triangle}_5\rangle \right).
\end{aligned}$$

Now reducing the pentagon:

$$\begin{aligned}
|\text{pentagon}\rangle &= e_1 \left(|\text{pentagon}_1\rangle + |\text{pentagon}_2\rangle + |\text{pentagon}_3\rangle + |\text{pentagon}_4\rangle + |\text{pentagon}_5\rangle \right) \\
&+ e_2 \left(|\text{pentagon}_6\rangle + |\text{pentagon}_7\rangle + |\text{pentagon}_8\rangle + |\text{pentagon}_9\rangle + |\text{pentagon}_{10}\rangle \right) \\
&= e_1 \left(d_1 \left(|\text{pentagon}_{11}\rangle + |\text{pentagon}_{12}\rangle \right) + d_2 \left(|\text{pentagon}_{13}\rangle + |\text{pentagon}_{14}\rangle \right) \right) + e_1 b |\text{pentagon}_{15}\rangle + e_1 c |\text{pentagon}_{16}\rangle \\
&+ e_1 c |\text{pentagon}_{17}\rangle + e_1 b |\text{pentagon}_{18}\rangle + e_2 \cdot 0 + e_2 |\text{pentagon}_{19}\rangle + e_2 b |\text{pentagon}_{20}\rangle + e_2 b |\text{pentagon}_{21}\rangle + e_2 |\text{pentagon}_{22}\rangle \\
&= (e_1 d_1 + e_1 b + e_1 c + e_2) |\text{pentagon}_{23}\rangle + (e_1 d_1 + e_1 c + e_1 b + e_2) |\text{pentagon}_{24}\rangle \\
&+ (e_1 d_2 + e_2 b) |\text{pentagon}_{25}\rangle + (e_1 d_2 + e_2 b) |\text{pentagon}_{26}\rangle
\end{aligned}$$

Hence getting the equations:

$$\begin{cases} e_1 d_1 + e_1 b + e_1 c + e_2 = c d_1 \\ e_1 d_2 + e_2 b = c d_2 \end{cases}$$

Next case is the square adjacent to the triangle. Reducing first the triangle:

$$|\hat{\triangle}\rangle = c |\triangle\rangle = c^2 |\text{triangle}\rangle$$

Now, reducing the square:

$$\begin{aligned} |\hat{\square}\rangle &= d_1 (|\triangle\rangle + |\hat{\triangle}\rangle) + d_2 (|\hat{\triangle}\rangle + |\triangle\rangle) \\ &= d_1 (c |\text{triangle}\rangle + b |\text{triangle}\rangle) + d_2 (0 + c |\text{triangle}\rangle) \\ &= (d_1 c + d_1 b + d_2 c) |\text{triangle}\rangle \end{aligned}$$

and we get the equations:

$$c^2 = d_1 c + d_1 b + d_2.$$

The next case is the square adjacent to a circle. In this case we shall represent the circle as a triangle, however, note that what differs this case from the previous one is that in this case there is no segment leaving the vertex of the triangle that is not adjacent to the square:

$$\begin{aligned} |\hat{\square}\rangle &= d_1 (|\hat{\triangle}\rangle + |\text{circle}\rangle) + d_2 (|\text{circle}\rangle + |\hat{\triangle}\rangle) \\ &= d_1 (0 + b |\text{circle}\rangle) + d_2 (|\text{circle}\rangle + a |\text{circle}\rangle) \\ &= (d_1 b + d_2 + a d_2) |\text{circle}\rangle \end{aligned}$$

Now, reducing the circle:

$$|\hat{\square}\rangle = b |\text{circle}\rangle = b^2 |\text{circle}\rangle.$$

and we get the equation:

$$b^2 = b d_1 + d_2 + a d_2.$$

The last case we must do is the pentagon adjacent to the circle, again we shall represent the circle as a triangle as we have done in the previous case. Reducing first the pentagon, we get:

$$\begin{aligned} |\hat{\pentagon}\rangle &= e_1 (|\triangle\rangle + |\hat{\triangle}\rangle + |\hat{\triangle}\rangle + |\hat{\triangle}\rangle + |\hat{\triangle}\rangle) \\ &\quad + e_2 (|\text{circle}\rangle + |\text{circle}\rangle + |\text{circle}\rangle + |\text{circle}\rangle + |\text{circle}\rangle) \\ &= e_1 c |\text{triangle}\rangle + e_1 \cdot 0 + e_1 b |\text{triangle}\rangle + e_1 b |\text{triangle}\rangle + e_1 \cdot 0 \\ &\quad + e_2 a |\text{triangle}\rangle + e_2 |\text{triangle}\rangle + e_2 \cdot 0 + e_2 \cdot 0 + e_2 |\text{triangle}\rangle \\ &= (e_1 c + 2e_1 b + 2e_2 + a e_2) |\text{triangle}\rangle \end{aligned}$$

and now reducing the circle:

$$\left| \text{circle with a dot} \right\rangle = b \left| \text{circle with a dot and a line} \right\rangle = bc \left| \text{circle with a dot and two lines} \right\rangle$$

hence, the equation:

$$bc = e_1c + 2e_1b + 2e_2 + ae_2.$$

Thus, to conclude, from all the previous relations, we get the set of equations:

$$\begin{cases} d_1e_2 + d_1d_2 = e_1e_2 + e_2c \\ d_1e_2 + d_2c = e_1e_2 + e_2b + 2e_1d_2 \\ d_1e_1 + d_2 = e_1^2 \\ d_1^2 + d_1e_1 = e_1^2 + e_1d_1 + e_1c \\ d_1e_2 = e_1e_2 + e_1d_2 \\ d_1e_1 = e_1^2 + e_1d_1 + e_2^2 \\ e_1d_1 + e_1b + e_1c + e_2 = cd_1 \\ e_1d_2 + e_2b = cd_2 \\ c^2 = d_1c + d_1b + d_2 \\ b^2 = bd_1 + d_2 + ad_2 \\ bc = e_1c + 2e_1b + 2e_2 + ae_2 \end{cases}$$

Now, having these equations we shall derive the values of each coefficient. By using the normalization $e_1 = 1$ (since $e_1 = 0$ produces only trivial solutions), we already get some simplifications, one of them allows us to easily compute e_2 :

$$d_1e_1 = e_1^2 + d_1e_1 + e_2 \Rightarrow e_1 = 1 + d_1 + e_2 \therefore e_2 = -1.$$

Now, let's name some of the equations above (after the simplification) just for future reference:

$$b^2 = bd_1 + d_2 + ad_2 \tag{1}$$

$$cd_1 = d_1 + c + b - 1 \tag{2}$$

$$d_1 + d_2 = 1 \tag{3}$$

$$d_1 + d_1^2 = 1 + c + d_1 \tag{4}$$

It is immediate From (3) that $d_2 = 1 - d_1$, as well as for (4) that $c = d_1^2 - 1$. Now, using those values of d_2 and c , in terms of d_1 in equation (2), we get:

$$\begin{aligned} (d_1^2 - 1)d_1 &= d_1 + (d_1^2 - 1) + b - 1 \\ \Rightarrow d_1^3 - d_1 &= d_1 + d_1^2 - 1 + b - 1 \\ \Rightarrow d_1^3 - d_1^2 - 2d_1 + 2 &= b \end{aligned}$$

And finally using this in equation (1):

$$\begin{aligned} (d_1^3 - d_1^2 - 2d_1 + 2)^2 &= (d_1^3 - d_1^2 - 2d_1 + 2)d_1 + 1 - d_1 + a(1 - d_1) \\ \Rightarrow d_1^6 - 2d_1^5 - 3d_1^4 + 8d_1^3 - 8d_1 + 4 &= d_1^4 - d_1^3 - 2d_1^2 + 2d_1 + 1 - d_1 + a(1 - d_1) \\ \Rightarrow d_1^6 - 2d_1^5 - 4d_1^4 + 9d_1^3 - 7d_1 + 3 &= a(1 - d_1) \end{aligned}$$

notice now, that since $1^6 - 2 \cdot 1^4 - 4 \cdot 1^4 + 9 \cdot 1^3 - 7 \cdot 1 + 3 = 0$, we can simplify further, and get the value of a in terms of d_1 as follows:

$$a = -5d_1^5 + d_1^4 + 5d_1^3 - 4d_1^2 - 4d_1 + 3.$$

Now, let's call $d_1 = -q - q^{-1}$, and then we get the exact equations established on the main theorem:

$$\begin{cases} a = q^5 + q^4 + 1 + 1 + q^{-1} + q^{-4} + q^{-5} \\ b = -q^3 - q^2 - q - q^{-1} - q^{-2} - q^{-3} \\ c = q^2 + 1 + q^{-2} \\ d_1 = -q - q^{-1} \\ d_2 = q + 1 + q^{-1} \\ e_1 = 1 \\ e_2 = -1 \end{cases}.$$

Now, we shall derive the expressions for the coefficients f_1, f_2, g_1 and g_2 . To do this, we will use the relations stated on the hypothesis of the the theorem together with the Reidemeister moves R_2 and the new one introduced for diagrams with triple points, the sliding over a vertex, which we shall denote by R_5 since we have already defined R_4 in the past. First, let's use this process on R_4 :

$$\begin{aligned}
\left| \text{Diagram 1} \right\rangle &= f_1 \left| \text{Diagram 2} \right\rangle + g_1 \left| \text{Diagram 3} \right\rangle + f_2 \left| \text{Diagram 4} \right\rangle + g_2 \left| \text{Diagram 5} \right\rangle \\
&= f_1 \left[f_1 \left| \text{Diagram 6} \right\rangle + g_1 \left| \text{Diagram 7} \right\rangle + f_2 \left| \text{Diagram 8} \right\rangle + g_2 \left| \text{Diagram 9} \right\rangle \right] \\
&\quad + g_1 \left[f_1 \left| \text{Diagram 10} \right\rangle + g_1 \left| \text{Diagram 11} \right\rangle + f_2 \left| \text{Diagram 12} \right\rangle + g_2 \left| \text{Diagram 13} \right\rangle \right] \\
&\quad + f_2 \left[f_1 \left| \text{Diagram 14} \right\rangle + g_1 \left| \text{Diagram 15} \right\rangle + f_2 \left| \text{Diagram 16} \right\rangle + g_2 \left| \text{Diagram 17} \right\rangle \right] \\
&\quad + g_2 \left[f_1 \left| \text{Diagram 18} \right\rangle + g_1 \left| \text{Diagram 19} \right\rangle + f_2 \left| \text{Diagram 20} \right\rangle + g_2 \left| \text{Diagram 21} \right\rangle \right] \\
&= f_1 \left[f_1 \left\{ d_1 \left| \text{Diagram 22} \right\rangle + d_1 \left| \text{Diagram 23} \right\rangle + d_2 \left| \text{Diagram 24} \right\rangle + d_2 \left| \text{Diagram 25} \right\rangle \right\} \right. \\
&\quad + g_1 \left\{ e_1 \left| \text{Diagram 26} \right\rangle + e_1 \left| \text{Diagram 27} \right\rangle + e_1 \left| \text{Diagram 28} \right\rangle + e_1 \left| \text{Diagram 29} \right\rangle + e_1 \left| \text{Diagram 30} \right\rangle \right\} \\
&\quad + g_1 \left\{ e_2 \left| \text{Diagram 31} \right\rangle + e_2 \left| \text{Diagram 32} \right\rangle + e_2 \left| \text{Diagram 33} \right\rangle + e_2 \left| \text{Diagram 34} \right\rangle + e_2 \left| \text{Diagram 35} \right\rangle \right\} + f_2 c \left| \text{Diagram 36} \right\rangle + g_2 \left| \text{Diagram 37} \right\rangle \Big] \\
&\quad + g_1 \left[f_1 c \left| \text{Diagram 38} \right\rangle + g_1 \left\{ d_1 \left| \text{Diagram 39} \right\rangle + d_1 \left| \text{Diagram 40} \right\rangle + d_2 \left| \text{Diagram 41} \right\rangle + d_2 \left| \text{Diagram 42} \right\rangle \right\} \right. \\
&\quad \left. + f_2 b \left| \text{Diagram 43} \right\rangle + g_2 \left| \text{Diagram 44} \right\rangle \right] \\
&\quad + f_2 \left[f_1 \left| \text{Diagram 45} \right\rangle + g_1 \left| \text{Diagram 46} \right\rangle + f_2 \left| \text{Diagram 47} \right\rangle + g_2 \left| \text{Diagram 48} \right\rangle \right] \\
&\quad + g_2 \left[f_1 b \left| \text{Diagram 49} \right\rangle + g_1 c \left| \text{Diagram 50} \right\rangle + 0 + g_2 \left| \text{Diagram 51} \right\rangle \right] \\
\Rightarrow \left| \text{Diagram 1} \right\rangle &= (f_1^2 d_1 + f_1 g_1 e_1 + g_1 f_1 c + g_1^2) \left| \text{Diagram 22} \right\rangle + (f_1^2 d_1 + f_1 g_1 e_1 + f_2 f_1) \left| \text{Diagram 23} \right\rangle \\
&\quad + (f_1 g_1 e_2 + f_1^2 d_2) \left| \text{Diagram 24} \right\rangle + (f_1^2 d_2 + f_1 g_1 e_2 + g_2 f_1 b + g_2 g_1 c + g_2^2) \left| \text{Diagram 25} \right\rangle \\
&\quad + (f_1 g_1 e_1 + f_1 g_2) \left| \text{Diagram 26} \right\rangle + (f_1 g_1 e_1 + g_1^2 d_1 + g_1 g_2) \left| \text{Diagram 27} \right\rangle \\
&\quad + (f_1 g_1 e_1 + f_2 g_1) \left| \text{Diagram 28} \right\rangle + (f_1 g_1 e_2 + f_2 g_2) \left| \text{Diagram 29} \right\rangle \\
&\quad + (f_1 g_1 e_2 + f_1 f_2 c + g_1^2 d_2 + g_1 f_2 b + f_2^2) \left| \text{Diagram 30} \right\rangle + (f_1 g_1 e_2 + g_1^2 d_2) \left| \text{Diagram 31} \right\rangle
\end{aligned}$$

On the other hand, the other side of the Reidemeister move R_5 , gives us:

$$\left| \text{Diagram 52} \right\rangle = f_1 \left| \text{Diagram 53} \right\rangle + g_1 \left| \text{Diagram 54} \right\rangle + f_2 \left| \text{Diagram 55} \right\rangle + g_2 \left| \text{Diagram 56} \right\rangle.$$

Hence, we get the equations:

$$\begin{cases} f_1^2 d_1 + f_1 g_1 e_1 + g_1 f_1 c + g_1^2 = 0 \\ f_1^2 d_1 + f_1 g_1 e_1 + f_2 f_1 = 0 \\ f_1 g_1 e_2 + f_1^2 d_2 = f_2 \\ f_1^2 d_2 + f_1 g_1 e_2 + g_2 f_1 b + g_2 g_1 c + g_2^2 = 0 \\ f_1 g_1 e_1 + f_1 g_2 = g_1 \\ f_1 g_1 e_1 + g_1^2 d_1 + g_1 g_2 = 0 \\ f_1 g_1 e_1 + f_2 g_1 = f_1 \\ f_1 g_1 e_2 + f_2 g_2 = 0 \\ f_1 g_1 e_2 + f_1 f_2 c + g_1^2 d_2 + g_1 f_2 b + f_2^2 = 0 \\ f_1 g_1 e_2 + g_1^2 d_2 = g_2 \end{cases}$$

If we do this same process but using the underpassing version of R_5 we get the same set of equations. Let's do for the R_2 move then:

$$\begin{aligned} |\curvearrowright\rangle &= f_1^2 |\text{X}\rangle + f_1 g_1 |\text{X}\rangle + f_1 f_2 |\text{X}\rangle + f_1 g_2 |\text{X}\rangle \\ &\quad + f_2 f_1 |\text{X}\rangle + f_2 g_1 |\text{X}\rangle + f_2^2 |\text{X}\rangle + f_2 g_2 |\text{X}\rangle \\ &\quad + g_1 f_1 |\text{X}\rangle + g_1^2 |\text{X}\rangle + g_1 f_2 |\text{X}\rangle + g_1 g_2 |\text{X}\rangle \\ &\quad + g_2 f_1 |\text{X}\rangle + g_2 g_1 |\text{X}\rangle + g_2 f_2 |\text{X}\rangle + g_2^2 |\text{X}\rangle \\ &= f_1^2 c |\text{X}\rangle + f_1 g_1 (d_1 |\text{X}\rangle + d_1 |\text{X}\rangle + d_2 |\text{X}\rangle + d_2 |\text{X}\rangle) + f_1 f_2 |\text{X}\rangle \\ &\quad + f_1 g_2 |\text{X}\rangle + f_2 f_1 |\text{X}\rangle + f_2 g_1 |\text{X}\rangle + f_2^2 |\text{X}\rangle + f_2 g_2 |\text{X}\rangle \\ &\quad + g_1 f_1 b |\text{X}\rangle + g_1^2 c |\text{X}\rangle + 0 + g_1 g_2 |\text{X}\rangle + 0 \\ &\quad + g_2 g_1 b |\text{X}\rangle + g_2 f_2 a |\text{X}\rangle + g_2^2 |\text{X}\rangle \\ \Rightarrow |\curvearrowright\rangle &= (f_1^2 c + f_1 g_1 d_1 + f_2 f_1 + g_1 f_1 b + g_1^2 c + g_1 g_2) |\text{X}\rangle + (f_1 g_1 d_1 + f_1 g_2 + f_2 g_1) |\text{X}\rangle \\ &\quad + (f_1 g_1 d_2 + f_1 f_2 b + f_1^2 + g_2 g_1 b + g_2 f_1 a + g_2^2) |\text{X}\rangle + (f_1 g_1 d_2 + f_2 g_2) |\text{X}\rangle \end{aligned}$$

And supposing this is invariant under the Reidemeister R_2 move, we get the following equations:

$$\begin{cases} f_1^2 c + f_1 g_1 d_1 + f_2 f_1 + g_1 f_1 b + g_1^2 c + g_1 g_2 = 0 \\ f_1 g_1 d_1 + f_1 g_2 + f_2 g_1 = 0 \\ f_1 g_1 d_2 + f_1 f_2 b + f_1^2 + g_2 g_1 b + g_2 f_1 a + g_2^2 = 0 \\ f_1 g_1 d_2 + f_2 g_2 = 1 \end{cases}.$$

Now, we shall solve for f_1, g_1, f_2 and g_2 the following equations:

$$\begin{aligned} f_1^2 c + f_1 g_1 d_1 + f_2 f_1 + g_1 f_1 b + g_1^2 c + g_1 g_2 &= 0 \\ f_1 g_1 d_1 + f_1 g_2 + f_2 g_1 &= 0 \\ f_1 g_1 d_2 + f_1 f_2 b + f_1^2 + g_2 g_1 b + g_2 f_1 a + g_2^2 &= 0 \\ f_1 g_1 d_2 + f_2 g_2 &= 1 \end{aligned} \tag{I}$$

$$\begin{aligned} f_1^2 d_1 + f_1 g_1 e_1 + g_1 f_1 c + g_1^2 &= 0 \\ f_1^2 d_1 + f_1 g_1 e_1 + f_2 f_1 &= 0 \\ f_1 g_1 e_2 + f_1^2 d_2 &= f_2 \\ f_1^2 d_2 + f_1 g_1 e_2 + g_2 f_1 b + g_2 g_1 c + g_2^2 &= 0 \end{aligned} \tag{II}$$

$$\begin{aligned} f_1 g_1 e_1 + f_1 g_2 &= g_1 \\ f_1 g_1 e_1 + g_1^2 d_1 + g_1 g_2 &= 0 \\ f_1 g_1 e_1 + f_2 g_1 &= f_1 \\ f_1 g_1 e_2 + f_2 g_2 &= 0 \end{aligned} \tag{III}$$

$$\begin{aligned} f_1 g_1 e_2 + f_1 f_2 c + g_1^2 d_2 + g_1 f_2 b + f_2^2 &= 0 \\ f_1 g_1 e_2 + g_1^2 d_2 &= g_2 \end{aligned} \tag{IV}$$

Remember that we already know that $e_1 = 1$ and $e_2 = -1$, by using this fact on the equations (I) and (III), we note that:

$$\begin{cases} f_1 g_1 d_2 + f_2 g_2 = 1 \\ -g_1 f_1 + g_2 f_2 = 0 \end{cases} \Rightarrow f_1 g_1 d_2 + g_1 f_1 = 1 \Rightarrow f_1 g_1 = \frac{1}{1 + d_2}$$

as well as $f_2 g_2 = \frac{1}{1 + d_2}$. Also, remember that $d_2 = q + 1 + q^{-1}$ and hence:

$$f_1 g_1 = f_2 g_2 = \frac{1}{q + 2 + q^{-1}}.$$

Now, consider the equations (II) and (IV). By multiplying them we get:

$$\begin{aligned} &-f_1^2 d_2 g_1 f_1 + f_1^2 d_2 g_1^2 + g_1^2 f_1^2 - g_1^2 d_2 g_1 f_1 = f_2 g_2 \\ \Rightarrow &(f_1^2 + g_1^2)(-d_2 g_1 f_1) + f_1^2 g_1^2 (d_2^2 + 1) = f_1 g_1 \\ \Rightarrow &(f_1^2 + g_1^2)(-d_2 g_1 f_1) = f_1 g_1 - (f_1 g_1)^2 (d_2^2 + 1) \\ \Rightarrow &f_1^2 + g_1^2 = \frac{f_1 g_1 [1 - f_1 g_1 (d_2^2 + 1)]}{-d_2 f_1 g_1} \\ \Rightarrow &f_1^2 + g_1^2 = \frac{1 - f_1 g_1 (d_2^2 + 1)}{-d_2} \\ \Rightarrow &f_1^2 + g_1^2 = \frac{1 - \frac{1}{1 + d_2} (d_2^2 + 1)}{-d_2} \\ \Rightarrow &f_1^2 + g_1^2 = - \left(\frac{1 + d_2 - (d_2^2 + 1)}{d_2 (1 + d_2)} \right) = - \left(\frac{d_2 - d_2^2}{d_2 + d_2^2} \right) \\ \therefore &f_1^2 + g_1^2 = \frac{d_2 - 1}{d_2 + 1} = \frac{q + q^{-1}}{2 + q + q^{-1}} \end{aligned}$$

Observe that, if we get $f_1 + g_1$ then the solution for f_1 and g_1 will be clear, since in this case we have the sum and product of two roots of a quadratic equation, to get this expression, we note that

$$(f_1 + g_1)^2 = f_1^2 + g_1^2 + 2f_1g_1$$

and since we know that $f_1^2 + g_1^2 = \frac{q + q^{-1}}{2 + q + q^{-1}}$ and $f_1g_1 = \frac{1}{q + 2 + q^{-1}}$, we easily obtain that:

$$(f_1 + g_1)^2 = \frac{q + q^{-1}}{q + 2 + q^{-1}} + \frac{2}{q + 2 + q^{-1}} = 1.$$

Therefore, taking $f_1 = \frac{1}{1 + q^{-1}}$ and $g_1 = \frac{1}{1 + q}$:

$$\frac{1}{1 + q^{-1}} + \frac{1}{1 + q} = \frac{(1 + q) + (1 + q^{-1})}{(1 + q)(1 + q^{-1})} = \frac{2 + q + q^{-1}}{1 + q + q^{-1} + 1} = 1$$

and

$$\left(\frac{1}{1 + q^{-1}} \right) \cdot \left(\frac{1}{1 + q} \right) = \frac{1}{q + 2 + q^{-1}}$$

hence, these are the solutions for f_1 and g_1 .

Now, to complete the proof we just need to verify the values of g_2 and f_2 , but by looking at equations (II) and (IV) again, we notice that, for finding the values of f_2 and g_2 we must know the values of f_1 , d_2 , e_2 and g_1 , which we already do and therefore substituting these known values on these equations we finally get:

$$f_2 = \frac{q}{1 + q^{-1}} \text{ and } g_2 = \frac{q^{-1}}{1 + q}.$$

□

Some of the calculations on the previous proof may be very confusing sometimes, thus we introduce now a simple way to do these calculations by hand. The main idea uses the definition of boundary of a planar graph as well as the definition of the endpoints. To explain this method we will use the reducing of the diagram of the left hand side of the Reidemeister move that we called R_5 ,  by applying the diagram .

We start by fixing the endpoints of the diagram (which is shown in red in the figure below) and drawing the first diagram. Now, we copy the endpoints, choose a crossing (represented in blue), and apply the reducing movement to that crossing, the whole diagram will be intact, except for that small area where we are applying the reducing movement.

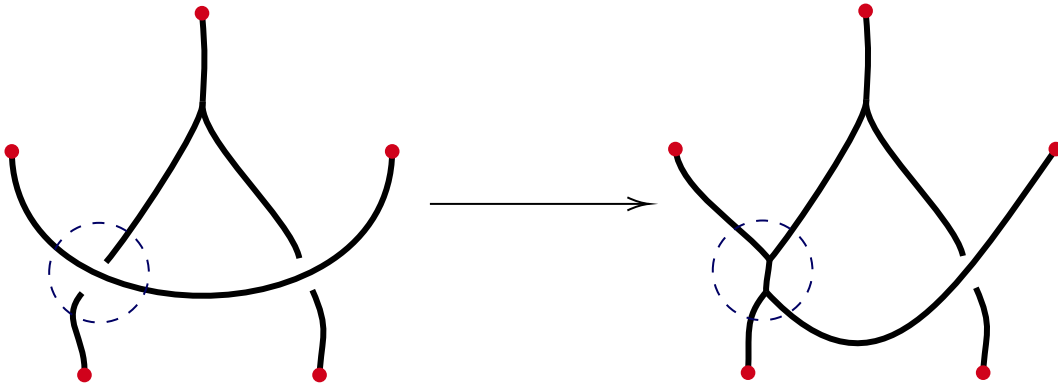


Figure 5.2: A method for calculating with diagrams.

By applying this method in every diagram we need, there will be no confusion in the calculations and even though they look confusing and messy, the results will easily follow.

6 Lie Algebras

6.1 The Definition and The Classical Lie Algebras

The main objective of this chapter is introduce concepts on Lie Algebras that may be important in future studies. Therefore, here we are more interested in understanding definitions and consequences of some results, instead of proving rigorously everything. But, since we are going to be following [11], all the proofs can be found in the same source ([11]).

Our goal, will be define and study the main results that leads us to formulate a theorem giving a classification for simple Lie algebras, in terms of Dyinkin diagrams, also introduce the classical and exceptional Lie algebras. Since here we are going to study Lie algebras, first we have to review what an algebra is:

Definition 6.1.1. Let $\mathbb{F}$ be a field and A a vector space over $\mathbb{K}$. We say that A is an **algebra** if there is a binary operation $\cdot: A \times A \rightarrow A$ satisfying the following conditions:

- (A₁) Right distributive: $(u + v) \cdot w = u \cdot w + v \cdot w$ for all $u, v, w \in A$;
- (A₂) Left distributive: $w \cdot (u + v) = w \cdot u + w \cdot v$ for all $u, v, w \in A$;
- (A₃) Compatibility with scalar multiplication: $(\alpha u) \cdot (\beta v) = (\alpha\beta)(u \cdot v)$ for all $u, v \in A$ and all $\alpha, \beta \in \mathbb{F}$.

We usually say that $\mathbb{F}$ is the **base field** of the algebra A .

A vector space L over a field $\mathbb{F}$, with an operation $[\cdot, \cdot]: L \times L \rightarrow L$ called **bracket**, is called a **Lie Algebra** over $\mathbb{F}$ if:

- (L₁) $[\cdot, \cdot]$ is bilinear;
- (L₂) $[x, x] = 0$, for all $x \in L$;
- (L₃) $[x, [y, z]] + [z, [x, y]] + [y, [z, x]] = 0$, for all $x, y, z \in L$ (Jacobi Identity).

From (L₁) and (L₂) we can obtain yet another "axiom" (which is not actually an axiom since it can be obtained from two axioms), that is $[x, y] = -[y, x]$.

Two Lie algebras L and L' are isomorphic if there exists $\varphi: L \rightarrow L'$ satisfying:

- φ is a vector-space isomorphism;
- $\varphi([x, y]) = [\varphi(x), \varphi(y)]$.

The set of all linear transformations from V to itself is denoted by $\text{End}(V)$, and by defining a bracket on $\text{End}(V)$ as $[f, g] = f \circ g - g \circ f$, not only $\text{End}(V)$ is a vector space over $\mathbb{F}$, but it is also a Lie Algebra. In order to distinguish the vector space $\text{End}(V)$ from the Lie algebra $(\text{End}(V), [\cdot, \cdot])$ we will denote the latter as $\mathfrak{gl}(V)$, and call it the **General Linear Algebra**.

We may identify $\mathfrak{gl}(V)$ with the set of $n \times n$ matrices, since the linear transformations between V and V may be seen as matrices when a basis is fixed on V . In this case, we denote $\mathfrak{gl}(V)$ as $\mathfrak{gl}(n, \mathbb{F})$.

6.1.1 The Classical Algebras: A_ℓ , B_ℓ , C_ℓ and D_ℓ

There are some important types of Lie algebras that we shall see, here we define them, and give some important information about them, such as dimension, but we shall not compute such things although their computation is not complicated and can be checked in [11].

The algebra A_ℓ : If we take V as a vector space with dimension $\ell + 1$, then we define the A_ℓ algebras to be the algebras:

$$\mathfrak{sl}(\ell + 1, \mathbb{F}) = \{\varphi \in \text{End}(V); \text{tr}(\varphi) = 0\}.$$

We call it the **Special Linear Algebra**. It's dimension is given by $\dim(\mathfrak{sl}(V)) = \ell + (\ell + 1)^2 - (\ell + 1)$.

An important Lie algebra that will appear throughout these notes is the algebra of type A_2 $\mathfrak{sl}(2, \mathbb{F}) = \mathfrak{sl}_2$, that is the algebra generated by the matrices:

$$x = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, y = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} \text{ and } h = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

The algebra C_ℓ : If we take V as a vector space with dimension 2ℓ and define f to be the non-degenerate skew-symmetric form f given by the matrix $s = \begin{pmatrix} 0 & I_\ell \\ -I_\ell & 0 \end{pmatrix}$, then we define the C_ℓ algebras to be the algebras:

$$\mathfrak{sp}(2\ell, \mathbb{F}) = \{\varphi \in \text{End}(V); f(\varphi(v), w) = -f(v, \varphi(w))\}.$$

We call it the **Symplectic Algebra**. It's dimension is given by $\dim(\mathfrak{sp}(V)) = 2\ell^2 + \ell$.

The algebra B_ℓ : If we take V as a vector space with dimension $2\ell + 1$ and define f to be the non-degenerate bilinear symmetric form f given by the matrix $s = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & I_\ell \\ 0 & I_\ell & 0 \end{pmatrix}$, then we define the C_ℓ algebras to be the algebras:

$$\mathfrak{o}(2\ell + 1, \mathbb{F}) = \{\varphi \in \text{End}(V); f(\varphi(v), w) = -f(v, \varphi(w))\}.$$

We call it the **Orthogonal Algebra**. It's dimension is given by $\dim(\mathfrak{sp}(V)) = 2\ell^2 + \ell$.

The algebra D_ℓ : If we take V as a vector space with dimension 2ℓ and define f to be the non-degenerate skew-symmetric form f given by the matrix $s = \begin{pmatrix} 0 & I_\ell \\ I_\ell & 0 \end{pmatrix}$, then we define the C_ℓ algebras to be the algebras:

$$\mathfrak{o}(2\ell, \mathbb{F}) = \{\varphi \in \text{End}(V); f(\varphi(v), w) = -f(v, \varphi(w))\}.$$

We call it the **Orthogonal Algebra**. It's dimension is given by $\dim(\mathfrak{sp}(V)) = 2\ell^2 - \ell$.

To distinguish the two orthogonal algebras (if necessary) we shall call B_ℓ the odd orthogonal algebra and D_ℓ the even orthogonal algebra.

6.1.2 Representations of Lie Algebras

The notions of homomorphisms, epimorphisms, monomorphisms, isomorphisms and kernels of Lie algebras arise as it does in the study of other algebraic structures. A *Representation* of a Lie algebra is a homomorphism:

$$\varphi: L \rightarrow \mathfrak{gl}(V)$$

V being a vector space over $\mathbb{F}$.

For example, we have the *adjoint representation* given by:

$$\begin{aligned} \text{ad}: L &\rightarrow \mathfrak{gl}(L) \\ x &\mapsto \text{ad}(x) \end{aligned}$$

where $\text{ad}(x)(y) = [x, y]$. Recall that $\mathfrak{gl}(L)$ is a space of linear transformations, therefore $\text{ad}(x)$ is a linear transformation. Note also that $\ker(\text{ad}) = Z(L)$, since for every $x \in \ker(\text{ad})$ we have $[x, y] = 0$ for every $y \in L$, therefore, $x \in Z(L)$, and the converse follows in the same way.

The consequence of this fact is that if we take L simple, i.e., it's ideals are only $\{0\}$ and L itself, then this implies that $Z(L) = \{0\}$, i.e., $\ker(\text{ad}) = \{0\}$, and hence ad is an isomorphism onto its image, therefore we get that any simple Lie algebra is isomorphic to a linear Lie algebra.

6.2 Classification of Semissimple Lie Algebras

6.2.1 Nilpotent Lie Algebras

Define a sequence of ideals of L given by $L^0 = L$, $L^1 = [L, L]$, $L^2 = [L, L^1]$, $\dots$, $L^i = [L, L^{i-1}]$ called the **descending (or lower) central series** of L . Here $[L, L]$ is the following set:

$$[L, L] = \{[u, v]; u, v \in L\}.$$

We will say that L is **nilpotent** if $L^n = 0$ for some $n \in \mathbb{Z}_+$. For example, if L is abelian then $[L, L] = 0$, hence $L^1 = 0$, therefore any abelian Lie algebra is nilpotent.

If we define the **center** of L as the set $Z(L) = \{u \in L; [u, v] = 0 \forall v \in L\}$, then we have the following facts about a Lie algebra L :

- (1) If L is nilpotent, then so are all its subalgebras and homomorphic images of L ;

- (2) If $\frac{L}{Z(L)}$ is nilpotent so is L ;
- (3) If L is nilpotent and non-zero, then $Z(L) \neq \{0\}$.

Now, defining the sequence $L^{(0)} = L$, $L^{(1)} = [L, L]$, $L^{(2)} = [L^{(1)}, L^{(1)}]$, $\dots$, $L^{(i)} = [L^{(i-1)}, L^{(i-1)}]$, we say that L is a **solvable** Lie algebra if there exists $n \in \mathbb{Z}_+$ such that $L^{(n)} = 0$. Then, we have the following theorem due to Lie:

Theorem 6.2.1 (Lie's Theorem). *Let L be a solvable subalgebra of $\mathfrak{gl}(V)$, $\dim(V) = n < \infty$. Then the matrices of L relative to a suitable basis of V are upper triangular.*

Proposition 6.2.1. *Let V be a finite dimensional vector space over $\mathbb{F}$, $x \in \text{End}(V)$:*

- (a) *There exist unique $x_s, x_n \in \text{End}(V)$ satisfying: $x = x_s + x_n$, where x_s is semisimple, x_n is nilpotent and x_s and x_n commute;*
- (b) *There exist polynomials $p(T)$ and $q(T)$ in one indeterminate, without constant term, such that $x_s = p(x)$, $x_n = q(x)$, x_s and x_n commute with an endomorphism commuting with x ;*
- (c) *If $A \subset B \subset V$ are subspaces, and x maps B into A , then x_s and x_n also map B into A .*

The decomposition $x = x_s + x_n$ is called **Jordan-Chevalley decomposition**. To say that x_s is **semisimple** is to say that the roots of its minimal polynomial are all distinct. To say that x_n is **nilpotent**, is to say that $x_n^k = \text{Id}$ for some $k \in \mathbb{N}$.

Let us briefly recall that an algebra is said to be **semisimple** if it can be factor into the direct sum of **simple** subalgebras, i.e., subalgebras that, as an algebra itself, have no non-trivial ideals, which means that the only ideals of each factor of a semisimple algebra are $\{0\}$ and the factor itself.

If for a Lie algebra L we define a bilinear form $\kappa(x, y) = \text{tr}(\text{ad}(x), \text{ad}(y))$ called the **Killing form**, then the following theorem gives a characterization (as well as a criterion) for the Lie algebra to be semisimple:

Theorem 6.2.2. *Let L be a Lie algebra. Then L is semisimple if, and only if, its Killing form is non-degenerate.*

Recall that a bilinear form B is non-degenerate when $B(x, \bullet)$ and $B(\bullet, y)$ are the identically zero forms only when $x = 0$ and $y = 0$ respectively.

6.2.2 Casimir Element of a Representation

Let $\varphi: L \rightarrow \mathfrak{gl}(V)$ be a **faithful** representation (i.e., a representation that is injective) with non-degenerate trace form $\beta(x, y) = \text{tr}(\varphi(x)\varphi(y))$, recall that a trace form is nothing more than a bilinear form satisfying $\beta(xy, z) = \beta(x, yz)$. In this case, having fixed a basis $(x_1, \dots, x_n)$ for L , we write c_φ as

$$c_\varphi(\beta) = \sum_i \varphi(x_i)\varphi(y_i) \in \text{End}(V)$$

where $(y_0, \dots, y_n)$ is the dual basis of the previous fixed basis, and call it the **Casimir element of φ** . Note that c_φ is an endomorphism of V commuting with $\varphi(L)$. In fact, if $x \in L$, then:

$$[x, x_i] = \sum_j a_{ij} x_j \text{ and } [x, y_i] = \sum_j b_{ij} y_j$$

then

$$\begin{aligned} a_{ik} &= \sum_j a_{ij} \beta(x_i, y_k) = \beta([x, x_i], y_k) = \beta(-[x_i, x], y_k) \\ &= \beta(x_i, -[x, y_k]) = -\sum_j b_{kj} \beta(x_i, y_j) = -b_{ki} \end{aligned}$$

Now,

$$\begin{aligned} [\varphi(x), c_\varphi(\beta)] &= \sum_i [\varphi(x), \varphi(x_i)] \varphi(y_i) + \sum_i \varphi(x_i) [\varphi(x), \varphi(y_i)] \\ &= \sum_{ij} a_{ij} \varphi(x_j) \varphi(y_i) + \sum_{ij} b_{ij} \varphi(x_i) \varphi(y_j) = 0 \text{ since } a_{ij} = -b_{ji} \end{aligned}$$

Example. Let $L = \mathfrak{sl}(2, \mathbb{F})$, $V = \mathbb{F}^2$. Take $\varphi: L \rightarrow \mathfrak{gl}(V)$ the identity map. Let (x, h, y) be the standard basis for L . Then the dual basis relative to the trace form is $\left(y, \frac{h}{2}, x\right)$, then $c_\varphi = xy + \frac{1}{2}h^2 + yx$, with matrix $\begin{pmatrix} \frac{3}{2} & 0 \\ 0 & \frac{3}{2} \end{pmatrix}$. Notice that $\frac{3}{2} = \frac{\dim(L)}{\dim(V)}$.

6.2.3 Toral Subalgebras

Note that if L is not nilpotent, there must exist some element $x \in L$ such that x is not nilpotent, in other words, in the Jordan decomposition $x = x_s + x_n$ we must have $x_s \neq 0$. Hence, L has non zero subalgebras, for example, that generated by the elements x_s for all $x \in L$. We call these **toral subalgebras**.

Lemma 6.2.1. *A toral subalgebra of L is abelian.*

Proof. Let T be toral. Since $\text{ad } x$ is semisimple and $\mathbb{F}$ is algebraically closed, then $\text{ad } x$ is diagonalizable. Thus we need to show that $\text{ad}_T x$ has no nonzero eigenvalues, so that $\text{ad}_T x = 0$ for all $x \in T$ which means that T is abelian.

Suppose $[x, y] = ay$ for some $a \neq 0$ and some $y \neq 0$ in T . Then $\text{ad}_T y(x) = -ay$ is an eigenvector of $\text{ad}_T y$ of eigenvalue 0, since $\text{ad}_T y(\text{ad}_T y) = \text{ad}_T y(-ay) = -a \text{ad}_T y(y) = -a \cdot 0 = 0$. On the other hand, since $y \in T$, then we can write x as a linear combination of eigenvalues of $\text{ad}_T y$. Applying $\text{ad}_T y$ to x , we get a linear combination of eigenvectors which belong to eigenvalues zero, This contradicts the supposition. $\square$

Now fix H a maximal toral algebra, meaning that H is not contained in any toral subalgebra of L). From the previous lemma, H is abelian, hence $\text{ad}_L H$ is a commuting family of semisimple endomorphisms of L , thus we may write:

$$L = \bigoplus_{\alpha \in H^*} L_\alpha; \quad L_\alpha = \{x \in L; [h, x] = \alpha(h)x, \forall h \in H\}$$

since from linear algebra, we have that L is simultaneously diagonalizable (if A and B are commuting, then they are simultaneously diagonalizable). The non-zero $\alpha \in H^*$ are

called **roots** and the corresponding spaces L_α are called **root spaces**. The above decomposition $L = \bigoplus L_\alpha$ is called **Cartan decomposition**.

Example. For $L = \mathfrak{sl}(2, \mathbb{F})$, a maximal toral subalgebra H is spanned by $h = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$. The roots are linear forms $\alpha_1: H \rightarrow \mathbb{F}$ such that $\alpha_1(h) = 2$ and $\alpha_2: H \rightarrow \mathbb{F}$ such that $\alpha_2(h) = -2$, since:

$$\begin{aligned} L_1 &= \{x \in L; [h, x] = \alpha_1(h)x = 2x\} = \langle x \rangle \\ L_2 &= \{x \in L; [h, x] = \alpha_2(h)x = -2x\} = \langle -x \rangle = \langle y \rangle. \end{aligned}$$

Proposition 6.2.2. *The following assertions are equivalent:*

- (i) $[L_\alpha, L_\beta] \subset L_{\alpha+\beta}$ for all $\alpha, \beta \in H^*$;
- (ii) $\text{ad } x$ is nilpotent for all $x \in L_\alpha$, $\alpha \neq 0$;
- (iii) $\kappa(L_\alpha, L_\beta) = 0$ for all $\alpha, \beta \in H^*$, $\alpha + \beta = 0$.

Proof. (i) $x \in L_\alpha$, $y \in L_\beta$, $h \in H$, then we have:

$$\begin{aligned} \text{ad } h([x, y]) &= [[h, x], y] + [x, [h, y]] \\ &= \alpha(h)[x, y] + \beta(h)[x, y] \\ &= (\alpha + \beta)(h)[x, y]. \end{aligned}$$

(ii) Follows directly from the previous item.

(iii) Take $h \in H$ such that $(\alpha + \beta)(h) \neq 0$. Then, if $x \in L_\alpha$, $y \in L_\beta$ since κ is associative, $\kappa([x, h], y) = -\kappa([x, h], y) = -\kappa(x, [h, y]) \Leftrightarrow \alpha(h)\kappa(x, y) = -\beta(h)\kappa(x, y) \Leftrightarrow (\alpha + \beta)(h)\kappa(x, y) = 0$. Thus we get $\kappa(x, y) = 0$. □

Observe that $L_0 = \{x \in L; [x, h] = 0\} = C_L(H)$, the centralizer of H in L . We can write the Cartan decomposition of L as:

$$L = C_L(H) \oplus \bigoplus_{\alpha \in H^*} L_\alpha.$$

From now on, we shall use the following result, which will not be proven here:

"Let H be a maximal toral subalgebra of L . Then $H = C_L(H)$."

6.2.4 Root Systems

The restriction of κ to K is non degenerate, which allows us to identify H with H^* in the following way: to $\varphi \in H^*$ corresponds a unique element $t_\varphi \in H$ satisfying $\varphi(t_\varphi) = \kappa(t_\varphi, h)$ for all $h \in H$. In particular, Φ corresponds to the subset $\{t_\alpha; \alpha \in H\}$ of H . Here, the set Φ is the set of all elements of H^* that are non zero, and such that $L_\alpha \neq 0$ for all $\alpha \in \Phi$, i.e., Φ is the set of roots of L .

Proposition 6.2.3. *We have:*

- (a) Φ spans H^* ;

- (b) If $\alpha \in \Phi$, then $-\alpha \in \Phi$;
- (c) $\alpha \in \Phi$, $x \in L_\alpha$, $y \in L_{-\alpha}$. Then $[x, y] = \kappa(x, y)t_\alpha$;
- (d) $\alpha \in \Phi \Rightarrow [L_\alpha, L_{-\alpha}]$ is one dimensional with basis t_α ;
- (e) $\alpha(t_\alpha) = \kappa(t_\alpha, t_\alpha) \neq 0$ for all $\alpha \in \Phi$;
- (f) $\alpha \in \Phi$ and $x_\alpha \neq 0 \in L_\alpha$ then there exists $y_\alpha \in L_{-\alpha}$ such that $x_\alpha, y_\alpha, h_\alpha = [x_\alpha, y_\alpha]$ span a three dimensional simple subalgebra of L isomorphic to $\mathfrak{sl}_2(\mathbb{F})$ via: $x_\alpha \mapsto x$, $y_\alpha \mapsto y$ and $h_\alpha \mapsto h$, where x, y and h are the generators of $\mathfrak{sl}_2(\mathbb{F})$;
- (g) $h_\alpha = \frac{2t_\alpha}{\kappa(t_\alpha, t_\alpha)}$.

Proof. (a) If Φ fails to span H^* then there must exist an element $h \in H$ such that $\alpha(h) = 0$ for all $\alpha \in \Phi$. But this means that $[h, L_\alpha] = 0$ for all $\alpha \in \Phi$. Since $[h, H] = 0$ (since H is commutative), then $[h, L] = 0$ or $h \in Z(L)$, which is absurd because $Z(L) = \{0\}$ for L semisimple.

- (b) Let $\alpha \in \Phi$. If $-\alpha \notin \Phi$, then $L_{-\alpha} = \{0\}$, hence $\kappa(L_\alpha, L_\beta) = 0$ for all $\beta \in \Phi$ (since $-\alpha \notin \Phi$, there must not exist $\beta \in \Phi$ with $\alpha + \beta = 0$), hence $\kappa(L_\alpha, L) = 0$ contradicting the non degeneracy of κ .
- (c) Let $\alpha \in \Phi$, $x \in L_\alpha$ and $y \in L_{-\alpha}$. Let $h \in H$ arbitrary. The associativity of κ implies:

$$\begin{aligned} \kappa(h, [x, y]) &= \kappa([h, x], y) = \alpha(h)\kappa(x, y) \\ &= \kappa(t_\alpha, h)\kappa(x, y) \\ &= \kappa(\kappa(x, y)t_\alpha, h) \\ &= \kappa(h, \kappa(x, y)t_\alpha) \end{aligned}$$

and this says that H is orthogonal to $[x, y] - \kappa(x, y)t_\alpha$ forcing $[x, y] = \kappa(x, y)t_\alpha$, then the result follows from non degeneracy of κ on H .

- (d) Part (c) shows that t_α spans $[L_\alpha, L_{-\alpha}]$ since $[L_\alpha, L_{-\alpha}] \neq \{0\}$. Let $0 \neq x \in L_\alpha$. If $\kappa(x, L_{-\alpha}) = 0$ then $\kappa(x, L) = 0$ which is not true since κ is non degenerate. Therefore we can find $0 \neq y \in L_{-\alpha}$ for which $\kappa(x, y) \neq 0$. By the item (c), $[x, y] \neq 0 \Rightarrow \kappa(x, y)t_\alpha \neq 0$ hence $[L_\alpha, L_{-\alpha}] \neq \{0\}$.
- (e) Suppose $\alpha(t_\alpha) = 0$, so that $[t_\alpha, x] = 0 = [t_\alpha, y]$ for all $x \in L_\alpha$ and $y \in L_{-\alpha}$. As in (d) we can find such x, y satisfying $\kappa(x, y) \neq 0$. Modifying one or the other by a scalar, we may assume $\kappa(x, y) = 1$. Then $[x, y] = \kappa(x, y)t_\alpha = t_\alpha$.

It follows that the subspace S of L spanned by x, y, t_α is a three dimensional solvable subalgebra $S \simeq \text{ad}_L(S) \subset \mathfrak{gl}(L)$. In particular, $\text{ad}_L(s)$ is nilpotent for all $s \in [S, S]$ so $\text{ad}_L(t_\alpha)$ is both semisimple and nilpotent, i.e., $\text{ad}_L(t_\alpha) = 0$. This says that $t_\alpha \in Z(L) = \{0\}$, contrary to the choice of t_α .

- (f) Given $0 \neq x_\alpha \in L_\alpha$, find $y_\alpha \in L_{-\alpha}$ such that $\kappa(x_\alpha, y_\alpha) = \frac{2}{\kappa(t_\alpha, t_\alpha)}$, this is possible due to (e) and the fact that $\kappa(x_\alpha, L_{-\alpha}) \neq 0$. Set $h_\alpha = \frac{2t_\alpha}{\kappa(t_\alpha, t_\alpha)}$, then we get by

- (c) that $[x_\alpha, y_\alpha] = h_\alpha$. Moreover $[h_\alpha, x_\alpha] = \frac{2}{\alpha(t_\alpha)}[t_\alpha, x_\alpha] = \frac{2\alpha(t_\alpha)}{\alpha(t_\alpha)}x_\alpha = 2x_\alpha$. And similarly, we get $[h_\alpha, y_\alpha] = -2y_\alpha$. So x_α, y_α and h_α span the three dimensional subalgebra of L with same multiplication table as $\mathfrak{sl}_2(\mathbb{F})$.
- (g) Recall that t_α is given by $\kappa(t_\alpha, h) = \alpha(h)$ for $h \in H$. This shows that $t_\alpha = -t_{-\alpha}$ and in view of the way h_α has been defined, the assertion follows immediately. $\square$

Since the restriction to H of the killing form κ is non degenerate, we may transfer the form to H^* letting $(\gamma, \delta) = \kappa(t_\gamma, t_\delta)$ for all $\gamma, \delta \in H^*$.

We know that Φ spans H^* , hence we may take a basis $\alpha_1, \dots, \alpha_\ell$ a basis of H^* consisting of roots. If $\beta \in \Phi$, then $\beta = \sum_{i=1}^{\ell} c_i \alpha_i$, for $c_i \in \mathbb{F}$.

We now claim that $c_i \in \mathbb{Q}$. In fact, for each $j = 1, \dots, \ell$

$$(\beta, \alpha_j) = \sum_{i=1}^{\ell} c_i (\alpha_i, \alpha_j) \Rightarrow \frac{2(\beta, \alpha_j)}{(\alpha_j, \alpha_i)} = \sum_{i=1}^{\ell} \frac{2c_i (\alpha_i, \alpha_j)}{(\alpha_i, \alpha_j)}$$

this is a system of ℓ equalities in ℓ unknowns c_i with rational coefficients. Since $\{\alpha_i\}_{i=1}^{\ell}$ is a basis for H^* , and the form is non degenerate, then $((\alpha_i, \alpha_j))_{1 \leq i, j \leq \ell}$ is non singular, hence the system admits a unique solution over $\mathbb{Q}$, proving the claim.

Now we conclude that the space $E_{\mathbb{Q}}$ spanned by all the roots, has $\mathbb{Q}$ -dimension $\ell = \dim_{\mathbb{F}} H^*$. Furthermore, we have:

$$\begin{aligned} \lambda, \mu \in H^* &\Rightarrow (\lambda, \mu) = \kappa(t_\lambda, t_\mu) \\ &= \sum_{\alpha \in \Phi} \alpha(t_\alpha) \alpha(t_\mu) \\ &= \sum_{\alpha \in \Phi} (\alpha, \lambda) (\alpha, \mu). \end{aligned}$$

In particular for $\beta \in \Phi$, $(\beta, \beta) = \sum_{\alpha \in \Phi} (\alpha, \beta)^2 \Rightarrow \frac{1}{(\beta, \beta)} = \sum_{\alpha \in \Phi} \frac{(\alpha, \beta)^2}{(\beta, \beta)^2} \in \mathbb{Q}$. Hence, $(\beta, \beta) \in \mathbb{Q}$, thus all the inner products $(\alpha, \beta) \in \mathbb{Q}$ in $E_{\mathbb{Q}}$, so we obtain a non degenerate form in $E_{\mathbb{Q}}$. Further, for every $\lambda \in E_{\mathbb{Q}}$ we have $(\lambda, \lambda) = \sum_{\alpha \in \Phi} (\alpha, \lambda)^2 \geq 0$, therefore the form on $E_{\mathbb{Q}}$ is positive definite.

Now, let E be the space $E = \mathbb{Q} \otimes_{\mathbb{Q}} E_{\mathbb{Q}}$, i.e., the vector space obtained by extending the base field $\mathbb{Q}$ of $E_{\mathbb{Q}}$ to $\mathbb{R}$. The form extends canonically to E and $\dim_{\mathbb{R}} E = \ell$. Hence, we have proved the following theorem:

Theorem 6.2.3. *Consider L, H, Φ and E as constructed above, then we have:*

- (a) Φ spans E and 0 does not belong to Φ ;
- (b) If $\alpha \in \Phi$, then $-\alpha \in \Phi$, but no other scalar multiple of α is a root;
- (c) If $\alpha, \beta \in \Phi$, then $\beta - \frac{2(\beta, \alpha)}{(\alpha, \alpha)}\alpha \in \Phi$;
- (d) If $\alpha, \beta \in \Phi$, then $2\frac{(\beta, \alpha)}{(\alpha, \alpha)} \in \mathbb{Z}$.

Recall what a reflection on the Euclidian space is. We call Φ a **root system** in the real Euclidean space E , we therefore have created a correspondence between (L, H) and (Φ, E) .

Any non zero vector α of E determines a reflection σ_α with reflecting hyperplane $P_\alpha = \{\beta \in E; (\beta, \alpha) = 0\}$, the explicit formula is given by:

$$\sigma_\alpha(\beta) = \beta - \frac{2(\beta, \alpha)}{(\alpha, \alpha)}\alpha.$$

From now on, by simplicity, we shall denote $\langle \alpha, \beta \rangle = 2\frac{(\beta, \alpha)}{(\alpha, \alpha)}$.

Lemma 6.2.2. *Let Φ be a finite set which spans E . Suppose all the reflections σ_α for $\alpha \in \Phi$ leave Φ fixed. If $\sigma \in \text{GL}(E)$ also leaves fixed Φ , fixes pointwise a hyperplane P of E and sends non zero $\alpha \in \Phi$ to $-\alpha$, then $\phi = \phi_\alpha$ and the reflection plane $P = P_\alpha$.*

A subset Φ of the Euclidean space E is called a **root system** in E if the following axioms are satisfied:

- (R₁) Φ is finite, spans E , and does not contain 0;
- (R₂) If $\alpha \in \Phi$, then the only multiples of $\alpha \in \Phi$ are $\pm\alpha$;
- (R₃) If $\alpha \in \Phi$, the reflection σ_α leaves Φ fixed;
- (R₄) If $\alpha, \beta \in \Phi$, then $\langle \alpha, \beta \rangle$.

This justifies the name given previously to Φ constructed so far.

Let now Φ be a root system in E . Denote by $\mathcal{W}$ the subgroup of $\text{GL}(E)$ generated by the reflections σ_α for $\alpha \in \Phi$. By (R₃) $\mathcal{W}$ permutes the set Φ , which by (R₁) is finite and spans E . This allows us to identify $\mathcal{W}$ with the subgroup of the symmetric group on Φ , in particular $\mathcal{W}$ is finite. $\mathcal{W}$ is called the **Weyl group of Φ** .

Lemma 6.2.3. *Let Φ be a root system in E , with Weyl group $\mathcal{W}$. If $\sigma \in \text{GL}(E)$ leaves Φ fixed, then $\sigma\sigma_\alpha\sigma^{-1} = \sigma_{\sigma(\alpha)}$, and $\langle \beta, \alpha \rangle = \langle \sigma(\beta), \sigma(\alpha) \rangle$ for every $\alpha, \beta \in \Phi$.*

Proof. We have: $\sigma\sigma_\alpha\sigma^{-1}(\sigma(\beta)) = \sigma\sigma_\alpha(\beta) \in \Phi$ since $\sigma_\alpha(\beta) \in \Phi$. But since $\sigma_\alpha(\beta) = \beta - \langle \beta, \alpha \rangle \alpha$, this equals:

$$\sigma(\beta - \langle \beta, \alpha \rangle \alpha) = \sigma(\beta) - \langle \beta, \alpha \rangle \sigma(\alpha).$$

Since $\sigma(\beta)$ runs over Φ as β runs over Φ , we conclude that $\sigma\sigma_\alpha\sigma^{-1}$ leaves Φ fixed, while fixing pointwise the hyperplane $\sigma(P_\alpha)$ and sending $\sigma(\alpha)$ to $-\sigma(\alpha)$. Then, by the Lemma 6.2.2 we have $\sigma\sigma_\alpha\sigma^{-1} = \sigma_{\sigma(\alpha)}$. But now, comparing the equation above with the equation $\sigma_{\sigma(\alpha)}(\sigma(\beta)) = \sigma(\beta) - \langle \sigma(\beta), \sigma(\alpha) \rangle \sigma(\alpha)$, we get:

$$\langle \beta, \alpha \rangle = \langle \sigma(\beta), \sigma(\alpha) \rangle.$$

□

We can regard $\mathcal{W}$ as a subgroup of $\text{Aut } \Phi$, since an endomorphism of Φ is the same as an automorphism of E keeping Φ fixed by the constructed identification. Denote by $\alpha^\vee = \frac{2\alpha}{(\alpha, \alpha)}$ the **dual** or **inverse** of $\alpha \in \Phi$ and $\Phi^\vee = \{\alpha^\vee; \alpha \in \Phi\}$ the **dual** or **inverse** of Φ .

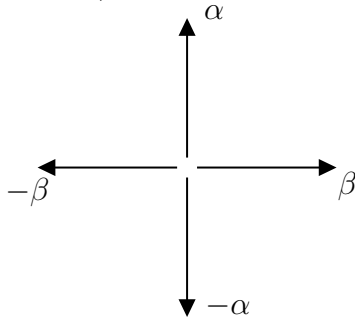
Examples. Call $\ell = \dim E$ the **rank** of the root system E . In view of (R_2) , there is only one possibility for the case $\ell = 1$:

$$-\alpha \longleftrightarrow \alpha \quad (A_1)$$

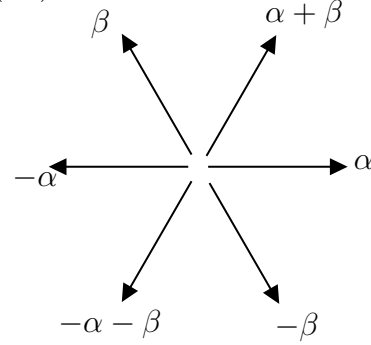
of course this is a root system with Weyl group of order 2, in Lie algebra, it belongs to $\mathfrak{sl}_2(\mathbb{F})$.

In rank 2, we have more possibilities:

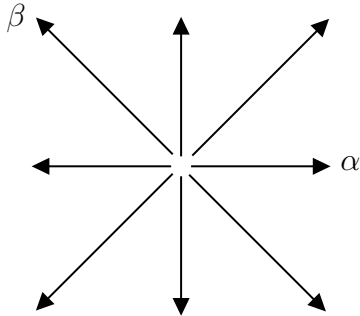
$(A_1 \times A_1)$



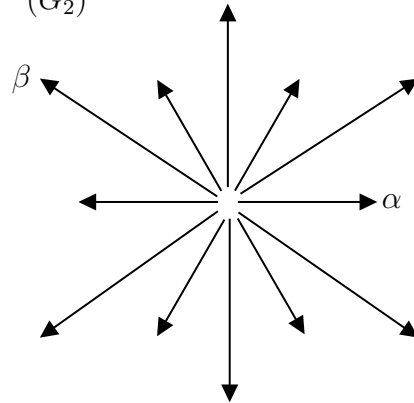
(A_2)



(B_2)



(G_2)



In view of axiom (R_4) , let's note that:

$$\langle \beta, \alpha \rangle = \frac{2(\beta, \alpha)}{(\alpha, \alpha)} = 2 \frac{\|\beta\|}{\|\alpha\|} \cos(\theta)$$

where θ is the angle between the vectors α and β . Also, $\langle \alpha, \beta \rangle \langle \beta, \alpha \rangle = 4 \cos^2(\theta) \in \mathbb{Z}_+$. But, $0 \leq \cos^2(\theta) \leq 1$ and $\langle \alpha, \beta \rangle \langle \beta, \alpha \rangle$ have the same sign as $\cos^2(\theta)$, so the following possibilities are the only ones when $\alpha \neq \pm \beta$ and $\|\beta\| \geq \|\alpha\|$:

$\langle \alpha, \beta \rangle$	$\langle \beta, \alpha \rangle$	θ	$\frac{\ \beta\ ^2}{\ \alpha\ ^2}$
0	0	$\frac{\pi}{2}$	undetermined
1	-1	$\frac{\pi}{3}$	1
-1	-1	$\frac{2\pi}{3}$	1
1	2	$\frac{\pi}{4}$	2
-1	-2	$\frac{3\pi}{4}$	2
1	3	$\frac{\pi}{6}$	3
-1	-3	$\frac{5\pi}{6}$	3

Lemma 6.2.4. *Let α and β be nonproportional roots. If $(\alpha, \beta) > 0$, then $\alpha - \beta$ is a root. If $(\alpha, \beta) < 0$, then $\alpha + \beta$ is a root.*

Proof. Since $(\alpha, \beta) > 0$, then $\langle \alpha, \beta \rangle > 0$, and from the Section 6.2.4, $\langle \alpha, \beta \rangle = 1$ or $\langle \beta, \alpha \rangle = 1$. If $\langle \alpha, \beta \rangle = 1$, then $\sigma_\beta(\alpha) = \alpha - \beta \in \Phi$ by using (R₃) and similarly, if $\langle \beta, \alpha \rangle = 1$, then $\beta - \alpha \in \Phi$ hence, $\sigma_{\beta-\alpha}(\beta - \alpha) = \alpha - \beta \in \Phi$. We prove the other case in a similar fashion. $\square$

Definition 6.2.1. Looking at all roots of the form $\beta + \alpha i$ for $i \in \mathbb{Z}$, we say that these elements form a **α -string through β** .

It can be shown that any root string is unbroken as well as their length is at most 4.

6.2.5 Cartan Matrix

A subset Δ of Φ is called a **base** if:

(B₁) Δ is a basis of E ;

(B₂) each root β can be written as $\beta = \sum_{\alpha \in \Delta} k_\alpha \alpha$ with $k_\alpha \in \mathbb{Z}$ for all $\alpha \in \Delta$, with k_α being all positive or all negative.

The roots in Δ are then called **simple roots**. We define the **height** of a root relative to Δ , to be the number $\text{ht}(\beta) = \sum_{\alpha \in \Delta} k_\alpha$.

Fix an ordering $(\alpha_1, \dots, \alpha_\ell)$ of the simple roots. The matrix $(\langle \alpha_i, \alpha_j \rangle)_{ij}$ is then called the **Cartan matrix** of Φ , and its entries are called **Cartan integers**.

Examples. $A_1 \times A_1 : \begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix}$, $A_2 : \begin{pmatrix} 2 & -1 \\ -1 & 2 \end{pmatrix}$, $B_2 : \begin{pmatrix} 2 & -2 \\ -1 & 2 \end{pmatrix}$, and $G_2 : \begin{pmatrix} 2 & -1 \\ -3 & 2 \end{pmatrix}$.

The interesting fact about the Cartan matrix is that it completely determines Φ .

Proposition 6.2.4. *Let $\Phi' \subset E'$ be another root system with basis $\Delta' = \{\alpha'_1, \dots, \alpha'_\ell\}$. If $\langle \alpha_i, \alpha_j \rangle = \langle \alpha'_i, \alpha'_j \rangle$ for $1 \leq i, j \leq \ell$, then the bijection $\alpha_i \mapsto \alpha'_i$ extends uniquely to an isomorphism $\varphi: E \rightarrow E'$ mapping Φ to Φ' and satisfying $\langle \varphi(\alpha), \varphi(\beta) \rangle = \langle \alpha, \beta \rangle$ for all $\alpha, \beta \in \Phi$. Therefore, the Cartan matrix of Φ determines Φ up to isomorphism.*

Proof. Since Δ is a basis for E , there is a unique vector space isomorphism sending $\varphi: E \rightarrow E'$ via $\alpha_i \mapsto \alpha'_i$. If $\alpha, \beta \in \Delta$, the hypothesis insures that

$$\begin{aligned} \sigma_{\varphi(\alpha)}(\varphi(\beta)) &= \sigma_{\alpha'}(\beta') = \beta' - \langle \beta', \alpha' \rangle \alpha' \\ &= \varphi(\beta) - \langle \beta, \alpha \rangle \varphi(\alpha) \\ &= \varphi(\beta - \langle \alpha, \alpha \rangle \alpha) = \varphi(\sigma_\alpha(\beta)) \end{aligned}$$

In other words we have commutative the following diagram:

$$\begin{array}{ccc} E & \xrightarrow{\varphi} & E' \\ \downarrow \sigma_\alpha & & \downarrow \sigma_{\varphi(\alpha)} \\ E & \xrightarrow{\varphi} & E' \end{array}$$

The respective Weyl groups $\mathcal{W}$ and $\mathcal{W}'$ are generated by the simple reflections, so it follows that the map $\sigma \mapsto \varphi \circ \sigma \circ \varphi^{-1}$ is an isomorphism of $\mathcal{W}$ onto $\mathcal{W}'$ sending σ_α to $\sigma_{\varphi(\alpha)}$ for $\alpha \in \Delta$.

But each $\beta \in \Phi$ is conjugate under $\mathcal{W}$ to a simple root, say $\beta = \sigma(\alpha)$ for $\alpha \in \Delta$. This implies that $\varphi(\beta) = (\varphi \circ \sigma \circ \varphi^{-1})(\varphi(\alpha)) \in \Phi'$. It follows that φ maps Φ to Φ' , moreover the formula for the reflection shows that φ preserves all Cartan integers. $\square$

The previous proof makes use of the following theorem, which will be stated without proof, since a proof would require some other concepts that, for now, will not be interesting for us:

Theorem 6.2.4. *Let Δ be a base of Φ :*

- (a) *If $\gamma \in E$, γ regular, there exists $\sigma \in \mathcal{W}$ such that $(\sigma(\gamma), \alpha) > 0$, for every $\alpha \in \Delta$;*
- (b) *If Δ' is another base of Φ then $\sigma(\Delta') = \Delta$ for some $\sigma \in \mathcal{W}$;*
- (c) *If α is any root, there exists $\sigma \in \mathcal{W}$ such that $\sigma(\alpha) \in \Delta$;*
- (d) *If $\sigma(\Delta) = \Delta$, for $\sigma \in \mathcal{W}$, then $\sigma = \text{Id}$.*

We now establish an algorithm for reconstructing Φ from Δ :

- Start by roots of height one (i.e., the simple roots);
- For any pair $\alpha_i \neq \alpha_j$, the integer r for some α_j -string through α_i is 0 (i.e., $\alpha_i - \alpha_j$ is not a root), so the integer q equals $-\langle \alpha_i, \alpha_j \rangle$. This enables us to write all the roots α of height 2, hence all the integers $\langle \alpha, \alpha_j \rangle$;
- For each root α of height 2, the integer r for the α_j -string through α can be determined easily, since α_j can be subtracted at most once, and then q is found because we know that $r - q = \langle \alpha, \alpha_j \rangle$.
- Repeat this process.

6.2.6 Dyinkin Diagrams

If α and β are distinct positive roots, then we know that $\langle \alpha, \beta \rangle \langle \beta, \alpha \rangle = 0, 1, 2$ or 3 . Define the i -th **Coxeter graph** of Φ to be a graph having ℓ vertices, the i -th joined to the j -th ($i \neq j$) by $\langle \alpha_i, \alpha_j \rangle \langle \alpha_j, \alpha_i \rangle$ edges. Examples:

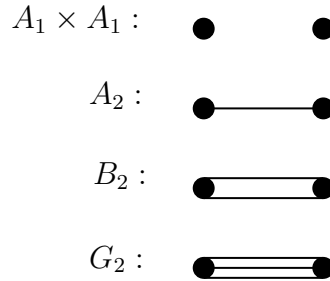


Figure 6.1: Coxeter graphs of $A_1 \times A_1$, A_2 , B_2 and G_2 .

The Coxeter graph determines the number $\langle \alpha_i, \alpha_j \rangle$ in case all roots have equal length, since $\langle \alpha_i, \alpha_j \rangle = \langle \alpha_j, \alpha_i \rangle$. In case more than one root length occurs, for example in the cases G_2 and B_2 , the graph fails to tell us which of a pair of vertices should correspond to a short simple root and which corresponds to a long simple root.

Whenever a double or a triple edge occurs in the Coxeter graph of Φ , we can add an arrow pointing to the shorter of the two roots. This additional information allows us to retrieve the Cartan integers, we call the resulting graph a **Dynkin diagram** of Φ . For example:

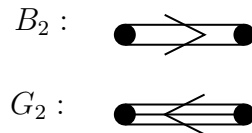
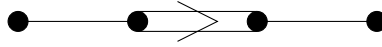


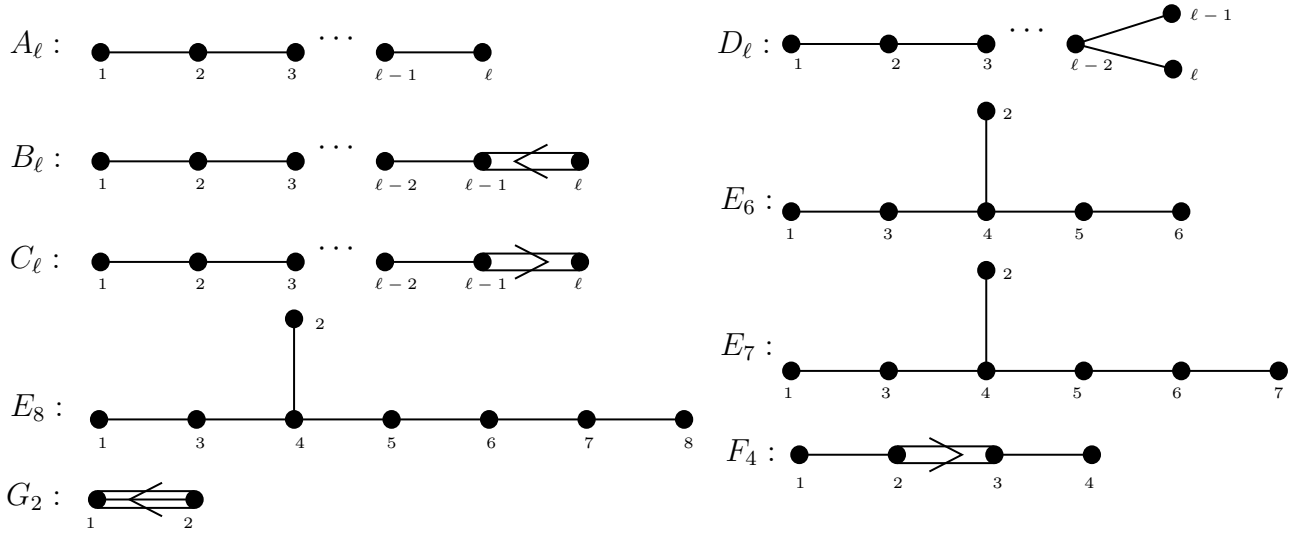
Figure 6.2: Dynkin diagrams of B_2 and G_2 .

Example. Given the diagram , we recover the Cartan matrix:

$$\begin{pmatrix} 2 & -1 & 0 & 0 \\ -1 & 2 & -2 & 0 \\ 0 & -1 & 2 & -1 \\ 0 & 0 & -1 & 2 \end{pmatrix}.$$

Note that non-diagonal entries of the Cartan matrix are always non-positive, since if they were positive, the angle between two simple roots would be $\frac{\pi}{3}$, which would imply that they may not be generators of Δ .

Theorem 6.2.5 (Classification Theorem). *If Φ is an irreducible root system of rank ℓ , its Dynkin diagram is one of the following:*



7 Topological Quantum Field Theories

This entire chapter will be closely following [22], hence, all the missing proofs can be found in the same source. This is, like the previous chapter, an introductory chapter, where our main goal is to achieve the definition of TQFT's preparing some previous knowledge for further studies in this area.

7.1 Monoidal Categories

We start by recalling that a category $\mathcal{V}$ is a collection of objects and arrows, called morphisms between these objects. To each morphism f there is two associated objects denoted by $\text{source}(f)$ and $\text{target}(f)$:

$$f: \text{source}(f) \rightarrow \text{target}(f).$$

For any pair of objects $V, W \in \mathcal{V}$, the set $\text{Hom}(V, W)$ denotes the set of all morphisms whose source is V and target is W . We also have the notion of a composition of morphisms with $\text{target}(g)=\text{source}(f)$, and this composition satisfied:

$$(f \circ g) \circ h = f \circ (g \circ h).$$

Finally, for each object there is an identity morphism $\text{Id}_V: V \rightarrow V$ such that $f \circ \text{Id}_V = f$ and $\text{Id}_V \circ g = g$ for every $f \in \text{Hom}(V, W)$ and $g \in \text{Hom}(W, V)$.

A *tensor product* in $\mathcal{V}$ is a covariant functor:

$$\otimes: \mathcal{V} \times \mathcal{V} \rightarrow \mathcal{V}$$

which associates to each pair of objects $V, W \in \mathcal{V}$ an object $V \otimes W \in \mathcal{V}$ and to each pair of morphisms $f: V \rightarrow V'$ and $g: W \rightarrow W'$ to a morphism $f \otimes g: V \otimes W \rightarrow V' \otimes W'$. The term covariant here, we remcall that means the following property is satisfied by the tensor product:

$$(f \circ f') \otimes (g \circ g') = (f \otimes g) \circ (f' \otimes g')$$

and

$$\text{Id}_V \otimes \text{Id}_W = \text{Id}_{V \otimes W}.$$

A *monoidal category* is a category $\mathcal{V}$ with a tensor product and an object $\mathbb{1} = \mathbb{1}_{\mathcal{V}}$ called the *unit* object, and such that the following properties hold:

- (1) $V \otimes \mathbb{1} = V, \mathbb{1} \otimes V = V$ for every object $V \in \mathcal{V}$;

- (2) $(U \otimes V) \otimes W = U \otimes (V \otimes W)$ for any objects $U, V, W \in \mathcal{V}$;
- (3) $f \otimes \text{Id}_1 = \text{Id}_1 \otimes f = f$ for any morphism $f \in \mathcal{V}$;
- (4) $(f \otimes g) \otimes h = f \otimes (g \otimes h)$ for any morphisms $f, g, h \in \mathcal{V}$.

the monoidal category shall be called *strict* if $V \otimes W = W \otimes V$, in general, $V \otimes W \simeq W \otimes V$ but not equal. We will freely use the term monoidal category referring to strict monoidal categories, unless it is necessary to distinguish them.

A very immediate example of a monoidal category is the category of modules over a commutative ring with the standard tensor product of modules. Though this category is not a strict monoidal category, meaning, $(U \otimes V) \otimes W \simeq U \otimes (V \otimes W)$ but not necessarily equal.

7.1.1 Braiding and Twists in Monoidal Categories

The previous example, gives us a monoidal category with a tensor product commutative up to isomorphism, meaning that there is an isomorphism between $V \otimes W$ and $W \otimes V$ taking $v \otimes w$ to $w \otimes v$, this isomorphism is called a *flip* and it will be denoted as $P_{V,W}$. For any three modules U, V and W we have the following properties (or commutativity of diagrams):

- $P_{U,V \otimes W} = (\text{Id}_V \otimes P_{U,W}) \circ (P_{U,V} \otimes \text{Id}_W)$.

$$\begin{array}{ccc}
 U \otimes V \otimes W & \xrightarrow{P_{U,V \otimes W}} & V \otimes W \otimes U \\
 \text{Id}_U \otimes P_{V,W} \downarrow & \nearrow \text{Id}_V \otimes P_{V,W} & \\
 V \otimes U \otimes W & &
 \end{array}$$

- $P_{U \otimes V, W} = (P_{U,W} \otimes \text{Id}_V) \circ (\text{Id}_U \otimes P_{V,W})$.

$$\begin{array}{ccc}
 U \otimes V \otimes W & \xrightarrow{P_{U \otimes V, W}} & B \\
 \text{Id}_U \otimes P_{V,W} \downarrow & \nearrow P_{U,W} \otimes \text{Id}_V & \\
 U \otimes W \otimes V & &
 \end{array}$$

Also, the set of all possible flips is involutive, i.e., $P_{W,V} \circ P_{V,W} = \text{Id}_{V \otimes W}$.

A *Braiding* in a monoidal category $\mathcal{V}$ consists of a natural family of isomorphisms

$$c = \{c_{V,W}: V \otimes W \rightarrow W \otimes V\}$$

where V and W are any objects in $\mathcal{V}$ and such that for any three objects U, V, W taken in $\mathcal{V}$ we have:

$$c_{U,V \otimes W} = (\text{Id}_V \otimes c_{U,W}) \circ (c_{U,V} \otimes \text{Id}_W).$$

Note also, that saying that this family is natural, is the same as saying that for any morphism $f: V \rightarrow V'$ and any morphism $g: W \rightarrow W'$ we have

$$(g \otimes f) \circ c_{V,W} = c_{V',W'} \circ (f \otimes g)$$

and in terms of diagram, it is to say that the following diagram commutes:

$$\begin{array}{ccc} V \otimes W & \xrightarrow{c_{V,W}} & W \otimes V \\ f \otimes g \downarrow & & \downarrow g \otimes f \\ V' \otimes W' & \xrightarrow{c_{V',W'}} & W' \otimes V' \end{array}$$

Also, $c_{V,1} = c_{1,V} = \text{Id}_V$ for any object $V \in \mathcal{V}$. Any braiding also satisfy the *Yang-Baxter identity* studied previously in the context of braiding and links, a very interesting connection of such an identity with the identity established earlier will appear when studying the ribbon graphs in the future:

$$(\text{Id}_W \otimes c_{U,V}) \circ (c_{U,W} \otimes \text{Id}_V) \circ (\text{Id}_U \otimes c_{V,W}) = (c_{V,W} \otimes \text{Id}_V) \circ (\text{Id}_V \otimes c_{U,W}) \circ (c_{U,V} \otimes \text{Id}_W)$$

diagrammatically:

$$\begin{array}{ccccc} & & U \otimes W \otimes V & \xrightarrow{c_{U,W} \otimes \text{Id}_V} & W \otimes U \otimes V \\ & \nearrow \text{Id}_U \otimes c_{V,W} & & & \searrow \text{Id}_W \otimes c_{U,V} \\ U \otimes V \otimes W & & & & W \otimes V \otimes U \\ & \searrow c_{U,V} \otimes \text{Id}_W & & & \nearrow c_{V,W} \otimes \text{Id}_U \\ & & V \otimes U \otimes W & \xrightarrow{\text{Id}_V \otimes c_{U,W}} & U \otimes W \otimes V \end{array}$$

Suppose $\mathcal{V}$ is a monoidal category with a braiding c , a *twist* in $\mathcal{V}$ consists of a natural family of isomorphisms:

$$\theta = \{\theta_V : V \rightarrow V\}$$

where V is any object in $\mathcal{V}$, and such that for any two objects V and W in the category $\mathcal{V}$, we have

$$\theta_{V \otimes W} = c_{W,V} \circ c_{V,W}(\theta_V \otimes \theta_W)$$

which is to say that the following diagram commutes:

$$\begin{array}{ccc} V \otimes W & \xrightarrow{\theta_{V \otimes W}} & V \otimes W \\ c_{V,W} \downarrow & & \uparrow \theta_V \otimes \theta_W \\ W \otimes V & \xrightarrow{c_{W,V}} & V \otimes W \end{array}$$

The natural condition means that for any morphism $f : U \rightarrow V$ in $\mathcal{V}$ we have $\theta_V \circ f = f \circ \theta_U$. Also, since c is also natural, we may rewrite the previous identity as

$$\theta_{V \otimes W} = c_{W,V} \circ (\theta_W \otimes \theta_V) \circ c_{V,W} = (\theta_V \otimes \theta_W) \circ c_{W,V} \circ c_{V,W}.$$

Note also that

$$(\theta_1)^2 = (\theta_1 \otimes \text{Id}_1) \circ (\text{Id}_1 \otimes \theta_1) = \theta_1 \otimes \theta_1 = \theta_1$$

which implies that $\theta_1 = \text{Id}_1$.

7.1.2 Duality in Monoidal Categories

Our objective on this chapter is to define the ribbon categories, and to finalize all the ingredients we need to construct such a category we still need to define a notion of duality in the monoidal category. Let $\mathcal{V}$ be a monoidal category. Assume that for each object V in $\mathcal{V}$ there exists an object V^* also in $\mathcal{V}$ and two morphisms:

$$b_V: \mathbb{1} \rightarrow V \otimes V^* \text{ and } d_V: V^* \otimes V \rightarrow \mathbb{1}$$

called, respectively, the *coevaluation* and the *evaluation* morphisms. Then we say that the rule $V \mapsto (V^*, b_V, d_V)$ is called a *duality* in $\mathcal{V}$ when the following two properties are satisfied:

$$\begin{aligned} (\text{Id}_V \otimes d_V) \circ (b_V \otimes \text{Id}_V) &= \text{Id}_V \\ (d_V \otimes \text{Id}_{V^*}) \circ (\text{Id}_{V^*} \otimes b_V) &= \text{Id}_{V^*}. \end{aligned}$$

We say that the duality on $\mathcal{V}$ is *compatible* with the braiding c and the twist θ in $\mathcal{V}$, if for any object taken in $\mathcal{V}$ we have

$$(\theta_V \otimes \text{Id}_{V^*}) \circ b_V = (\text{Id}_V \otimes \theta_{V^*}) \circ b_V.$$

We say that a duality compatible with braiding and twist is *involution* if $V^{**} = (V^*)^*$ is canonically isomorphic to V .

7.2 Ribbon Categories

Now, we have all the tools needed to define the ribbon categories, and our job from now on will be determine an explanation for such a name, after studying colored graphs we'll see how we can translate properties from ribbon categories to more visual contexts (the ribbon graphs). A *ribbon category* is a monoidal category $\mathcal{V}$ equipped with a braiding c , a twist θ and a duality $(*, b, d)$. strict monoidal categories give rise to *strict ribbon categories*, but we shall freely use only ribbon category when referring to the strict ones, unless it is necessary to distinguish them.

To each ribbon category $\mathcal{V}$ we can associate a mirror ribbon category $\bar{\mathcal{V}}$. This category is defined the same way $\mathcal{V}$ is, but the braiding $\bar{c}$ and the twist $\bar{\theta}$ in $\bar{\mathcal{V}}$ are defined as:

$$\bar{c}_{V,W} = (c_{W,V})^{-1} \text{ and } \bar{\theta}_V = (\theta_V)^{-1}$$

recall also that this mirror ribbon category is well defined, since both c and θ are families of isomorphisms, hence they have well defined inverses.

7.2.1 Traces and Dimensions

Let $\mathcal{V}$ be a ribbon category. We denote by $\mathbb{K} = \mathbb{K}_{\mathcal{V}}$ the semigroup $\text{End}(\mathbb{1})$. $\mathbb{K}$ is commutative since for any $k, k': \mathbb{1} \rightarrow \mathbb{1}$ we have:

$$k \circ k' = (k \otimes \text{Id}_{\mathbb{1}}) \circ (\text{Id}_{\mathbb{1}} \otimes k') = k \otimes k' = (\text{Id}_{\mathbb{1}} \otimes k') \circ (k \otimes \text{Id}_{\mathbb{1}}) = k' \otimes k.$$

The trace of morphisms and dimensions of objects of $\mathcal{V}$ will take values in $\mathbb{K}$. We may establish an analogy between traces and dimensions of morphisms and objects in ribbon categories to traces and dimensions of linear transformations and $\mathbb{F}$ -vector spaces,

where here the field $\mathbb{F}$ can be thought of as the semigroup $\mathbb{K}$ of the ribbon categories, the difference here is that the study of ribbon categories is much more general than the study of vector spaces.

For an endomorphism $f: V \rightarrow V$ of an object V we define the *trace* of f , denoted as $\text{tr}(f) \in \mathbb{K}$ as

$$\text{tr}(f) = f_V \circ c_{V,V^*} ((\theta \circ f) \otimes \text{Id}_{V^*}) \circ b_V \in \mathbb{K}.$$

For an object V in $\mathcal{V}$, we define its dimension, denoted as $\dim(V)$ as:

$$\dim(V) = \text{tr}(\text{Id}_V) = d_V \circ c_{V,V^*} \circ (\theta_V \otimes \text{Id}_{V^*}) \circ b_V \in \mathbb{K}.$$

The next lemma follows immediately from the definitions above:

Lemma 7.2.1. *We have:*

(i) *For any morphisms $f: V \rightarrow W$ and $g: W \rightarrow V$ we have:*

$$\text{tr}(f \circ g) = \text{tr}(g \circ f).$$

(ii) *For any endomorphisms f, g of objects in $\mathcal{V}$, we have:*

$$\text{tr}(f \otimes g) = \text{tr}(f) \circ \text{tr}(g).$$

(iii) *For any morphism $k: \mathbb{1} \rightarrow \mathbb{1}$ we have $\text{tr}(k) = k$.*

From this lemma we also get the following properties for dimensions, since dimension is defined in terms of trace:

(i)' $V \simeq W \Rightarrow \dim(V) = \dim(W)$;

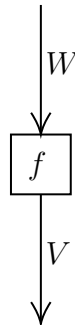
(ii)' For every V and W in $\mathcal{V}$ we have $\dim(V \otimes W) = \dim(V) \circ \dim(W)$;

(iii)' $\dim(\mathbb{1}) = 1$.

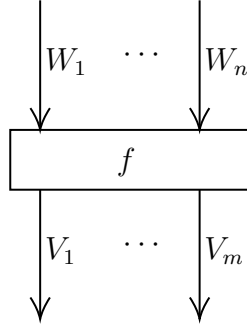
7.2.2 A Visual Study of Ribbon Categories

We now look for a way to visualize all the previous concepts, we do that by establishing some form of graphical calculus on the category $\mathcal{V}$. Suppose that $\mathcal{V}$ is a ribbon category and represent a morphism $f: V \rightarrow W$ by a box with two vertical arrows pointing downwards, we label the arrow going in the box W and the arrow going out of the box V , while the box is labeled f .

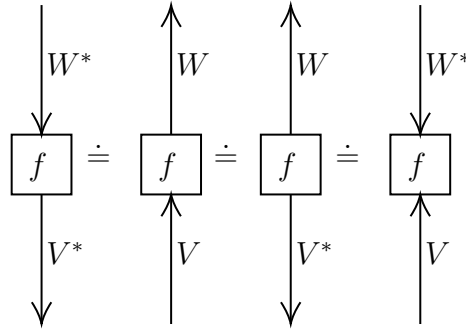
The labels we give to the arrows and boxes are called *colors* of such a graph. We get the following graph representing such f :



More generally, a morphism $f: V_1 \otimes \cdots \otimes V_m \rightarrow W_1 \otimes \cdots \otimes W_n$ may be represented the same way, but with n arrows going in the box, one for each term W_i and m arrows going out the box, one for each V_j term:

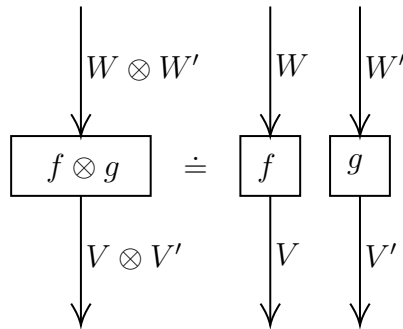


We may also represent the duals of each object as an arrow pointing in the opposite direction, for example, take $g: V^* \rightarrow W$, V and W as in the previous definition, since the arrow colored with V is pointing downwards (outside the box), the arrow colored with V^* should be pointing on the opposite direction, i.e., pointing towards the box. We get the following equality between graphs, and here we establish a new notation, every time we talk about equality between graphs we shall use the symbol $\doteq$:



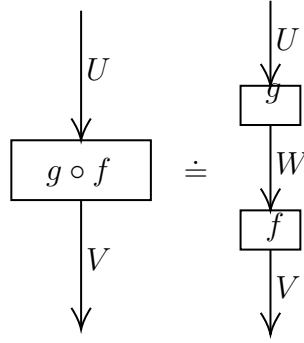
We shall represent the identity endomorphism on V by just an arrow pointing downwards with no box in it and colored with V . This means that every time we see an arrow pointing upwards colored with V we understand it as the identity endomorphism on V^* . Also, the empty picture represents the identity endomorphism on $\mathbb{1}$, since $V \otimes \mathbb{1} = V$.

We will also define the graph of the tensor product of two morphisms f and g as the graph of f put besides the graph of g , for example, if we consider $f: V \rightarrow W$ and $g: V' \rightarrow W'$, we have that $f \otimes g: V \otimes V' \rightarrow W \otimes W'$ is of the form:



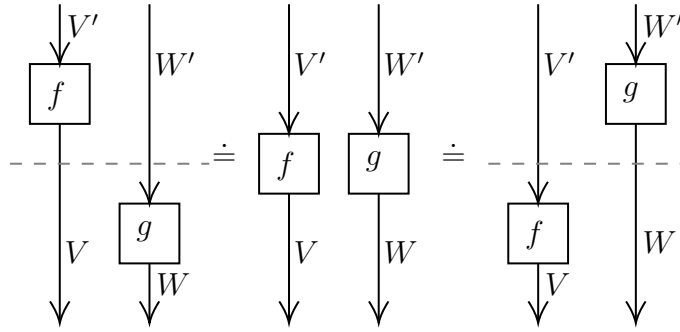
And finally, we define the composition of two morphisms $f: V \rightarrow W$ and $g: W \rightarrow U$

as being the graph of g put on top of the graph of f connecting the two corresponding arrows colored with W . Thus, we get:



Now, to exemplify the power of such a representation for the morphisms in the ribbon category $\mathcal{V}$, we will look visually the following equality for the morphisms $f: V \rightarrow V'$ and $g: W \rightarrow W'$:

$$(f \otimes \text{Id}_{W'}) \circ (\text{Id}_D \otimes g) = f \otimes g = (\text{Id}_{V'} \otimes g) \circ (f \otimes \text{Id}_W).$$

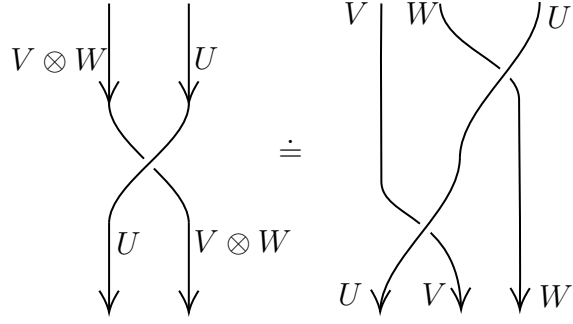


Now, we will see that by representing the properties of ribbon categories in a certain way on the level of graphs we can create some intuitive way of working with the morphisms just as if we were working with links in the past. Consider $c_{V,W}: V \otimes W \rightarrow W \otimes V$ a braiding, with inverse $c_{V,W}^{-1} = c_{W,V}: W \otimes V \rightarrow V \otimes W$, we represent it as:

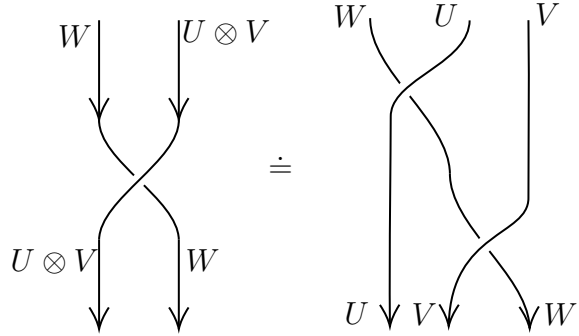


Therefore, we get the following pictures related to the properties of braiding established earlier:

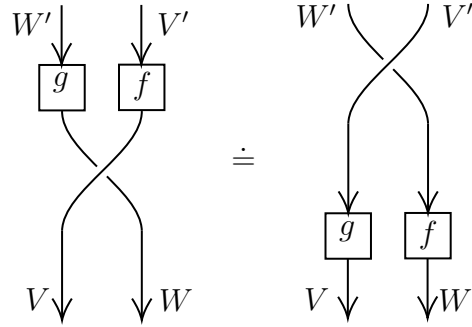
$$c_{U,V \otimes W} = (\text{Id}_V \otimes c_{U,W}) \circ (c_{U,V} \otimes \text{Id}_W):$$



$$c_{U \otimes V, W} = (c_{U,W} \text{Id}_V) \circ (\text{Id}_U \otimes c_{V,W}):$$

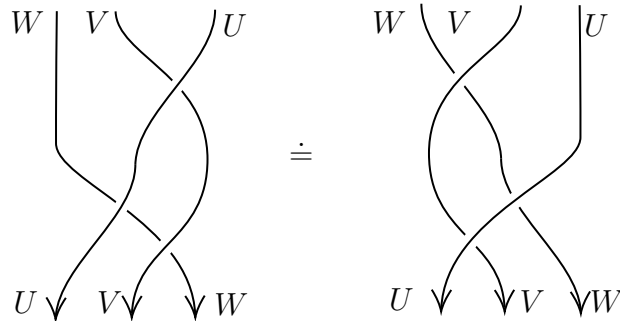


$$(g \otimes f) \circ c_{V,W} = c_{V',W'} \circ (f \otimes g):$$



We may also write the Yang-Baxter identity, and if we go back for a while, we see that this graph representation of the Yang-Baxter identity is compatible with that studied earlier on the context of braids:

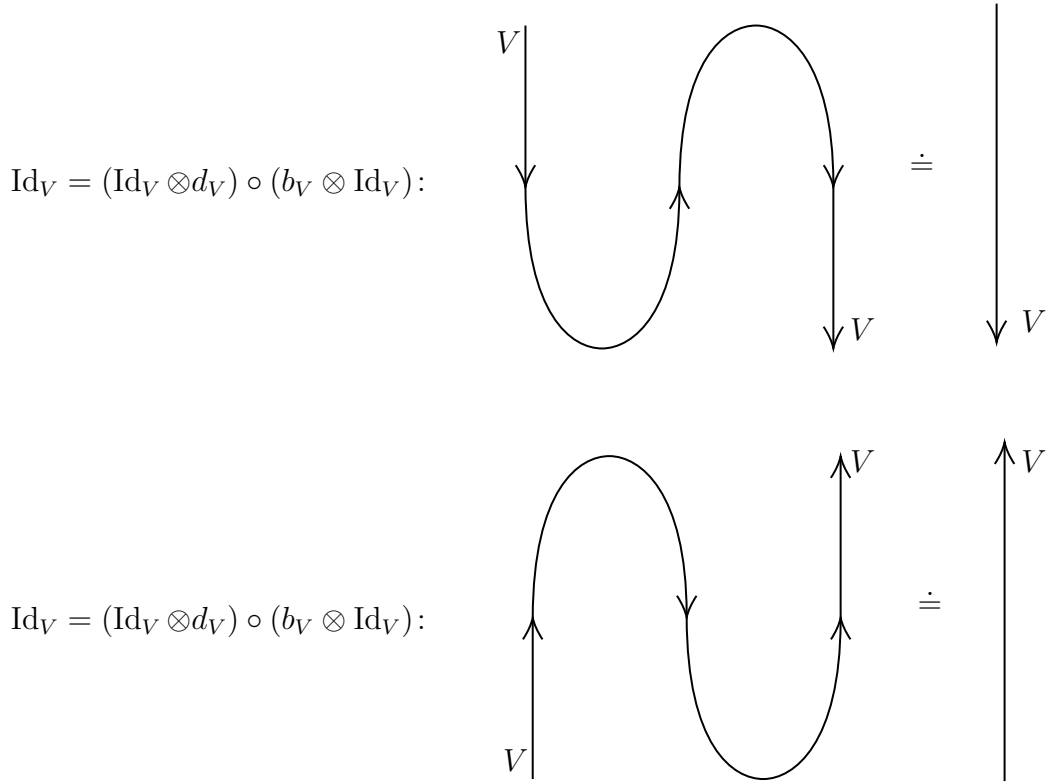
$$(\text{Id}_W \otimes c_{U,V}) \circ (c_{U,V} \otimes \text{Id}_V) \circ (\text{Id}_U \otimes c_{V,W}) = (c_{V,W} \otimes \text{Id}_V) \circ (\text{Id}_V \otimes c_{U,W}) \circ (c_{U,V} \otimes \text{Id}_W):$$



The dual morphisms, evaluation and coevaluation are represented by right oriented cap and cup respectively, since $\mathbb{1}$ is represented by the empty graph, we may imagine that taking the tensor product of two objects to $\mathbb{1}$ should result in the two arrows representing these objects colliding into a point, as represented in the next picture:



Finally, we may see graphically how the dual identities defined earlier are somewhat intuitive, as we may imagine it as just stretching some wobbled strand:



7.3 Ribbon Graph Category

Now that we've seen that we are able to represent the ribbon categories in a graphical fashion, which may remind us of braids, we need to understand if this set of graphs have some form of structure. And indeed, if we consider some properties on these graphs, such as coloring of arrows and boxes with certain objects and morphisms, respectively, on a certain monoidal category, we are able to construct a new category, which we shall call the *category of ribbon graphs*.

The importance of the construction of such a category, is that later we will see that we may construct a functor between the a monoidal category and a ribbon graph category, therefore providing the means we needed to work with objects and morphisms of monoidal categories, but in the contexts of ribbon gaphs, which may be easier to visualize and more intuitive to understand certain properties that may occur, examples of these were given in the previous section.

Definition 7.3.1. We define a *band* as the square $[0, 1] \times [0, 1]$ or a homeomorphic image of it. The images of $[0, 1] \times \{0\}$ and $[0, 1] \times \{1\}$ are called the *bases of the band*. The image of $\left\{\frac{1}{2}\right\} \times [0, 1]$ is called the *core of the band*.

Definition 7.3.2. An *annulus* is the cylinder $S^1 \times [0, 1]$, ou a homeomorphic image of this cylinder. The image of $\left\{\frac{1}{2}\right\} \times [0, 1]$ is called the *core of the annulus*.

A band or annulus is said to be *directed* if its core is oriented. The orientation of this core is called the *direction* of the band or the annulus. Making a connection with the previous section, the core of the oriented bands on a ribbon graph can be seen as the arrows on the diagrams we have drawn.

Definition 7.3.3. A *coupon* is a band with distinguished base. This distinguished base is called *bottom base* of the coupon and the opposite base is called the *top base* of the coupon.

Again, making the connection with the previous diagrams, the coupons are nothing but the boxes we have drawn. Note, however, that in the previous diagrams, we only drew the core of band, the whole band should be thought of as a kind of a framing of the drawn cores (or arrows).

Definition 7.3.4. Let $k, l \in \mathbb{Z}_+$. A *ribbon (k, l) -graaph* in $\mathbb{R}^3$ is an oriented surface Ω embedded in the strip $\mathbb{R}^2 \times [0, 1]$ and decomposed into a finite union of annuli, bands and coupons, such that:

- i) Ω meets the plane $\mathbb{R}^2 \times \{0\}$ and $\mathbb{R}^2 \times \{1\}$ orthogonally along the following segments:

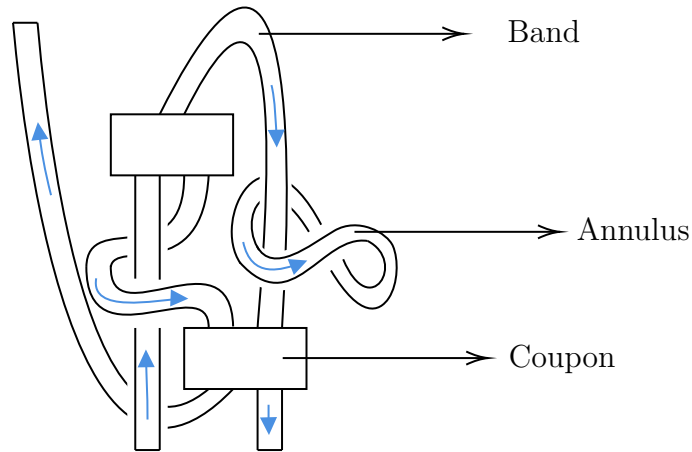
$$\left\{ \left[i - \left(\frac{1}{10} \right), i - \left(\frac{1}{10} \right) \right] \times \{0\} \times \{0\}; i = 1, \dots, k \right\},$$

$$\left\{ \left[j - \left(\frac{1}{10} \right), j - \left(\frac{1}{10} \right) \right] \times \{0\} \times \{1\}; j = 1, \dots, l \right\}$$

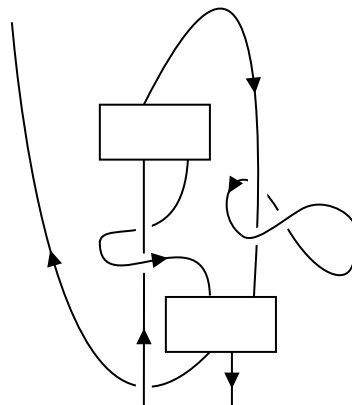
called the *bases* of certain bands of Ω ;

- ii) Other bases of the bands lie on the bases of coupons, otherwise the bands, coupons and annuli are disjoint;
- iii) The bands and annuli of Ω are directed.

The following picture represents such a ribbon graph:



Another way to represent these graphs is by representing only the core of bands and annuli, this is useful to simplify things, however if we are not careful we may lose some intuition when we deal with twists on the bands, therefore we need a way to represent such a thing on diagrams with no width. For the previous ribbon graph, for example, we have the following representation using only the core of the bands (note that the previous graph shows no twists on the bands):



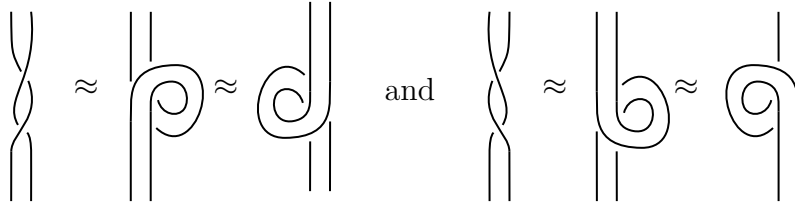
As we did before when working with braids, we also ask for a ribbon graph to be in *general position*, we define a graph to be in general position if it satisfy the following conditions:

- The graph should be very close and almost parallel to $\mathbb{R} \times \{0\} \times \mathbb{R}$;
- The coupons should be plane rectangles parallel to $\mathbb{R} \times \{0\} \times \mathbb{R}$;
- The bases of the coupons should be parallel to $\mathbb{R} \times \{0\} \times \{0\}$;
- The top base of a coupon should be higher than the bottom base;
- The projections of the arcs of bands and annuli in $\mathbb{R} \times \{0\} \times \mathbb{R}$ should have only double transversal crossings and should not overlap with the projections of the coupons.

In general, these requirements are technicalities and it shall nt be a problem for us on the studies, we will always consider the ribbon graphs in general position. The intuition we shall have on putting the graphs in general position is to present pictures of ribbon graphs by planar pictures that generalize our notion of knot diagrams.

Given a ribbon graph, after deforming it into a graph in general position, we may draw the projections of the cores of bands and annui taking into account the over and under crossing information. Finally, the resulting picture is what we call the *diagram of the ribbon graph*.

Note that given any ribbon graph, not necessarily in general position, we may deform it in such a way to get it to be in general position, in fact, first we deform it so that the coupons lie in general position, then we deform the bands so that they are parallel to the plane of projection. Now the only problem will happen when twists appear in the graph, we may deform the positive and negative twists in the following way:



We introduce now the coloring of bands, annuli and coupons of ribbon graphs, this is what will allow us to make a connection between the ribbon categories and the category of ribbon graphs. Consider $\mathcal{V}$ a strict monoidal category with duality. As before, we say that a ribbon graph is *colored* over $\mathcal{V}$ if each band and each annulus is equipped with an object of $\mathcal{V}$, the coupons of Ω may or may not be equipped with a morphism.

Consider Q a coupon of a ribbon graph Ω , $V_1, V_2, \dots, V_m$ the colors bands with bases on the bottom of Q and $W_1, W_2, \dots, W_n$ the colors of bands with bases on the top of Q , let $\varepsilon_1, \dots, \varepsilon_m \in \{-1, 1\}$ and $\nu_1, \dots, \nu_n \in \{-1, 1\}$ numbers determined by the direction of the bands in the following way:

- If the direction of the i -th band is outwards the coupon, then $\varepsilon_i = +1$;
- If the direction of the i -th band is inwards the coupon, then $\varepsilon_i = -1$.

then, a color for Q is a morphism:

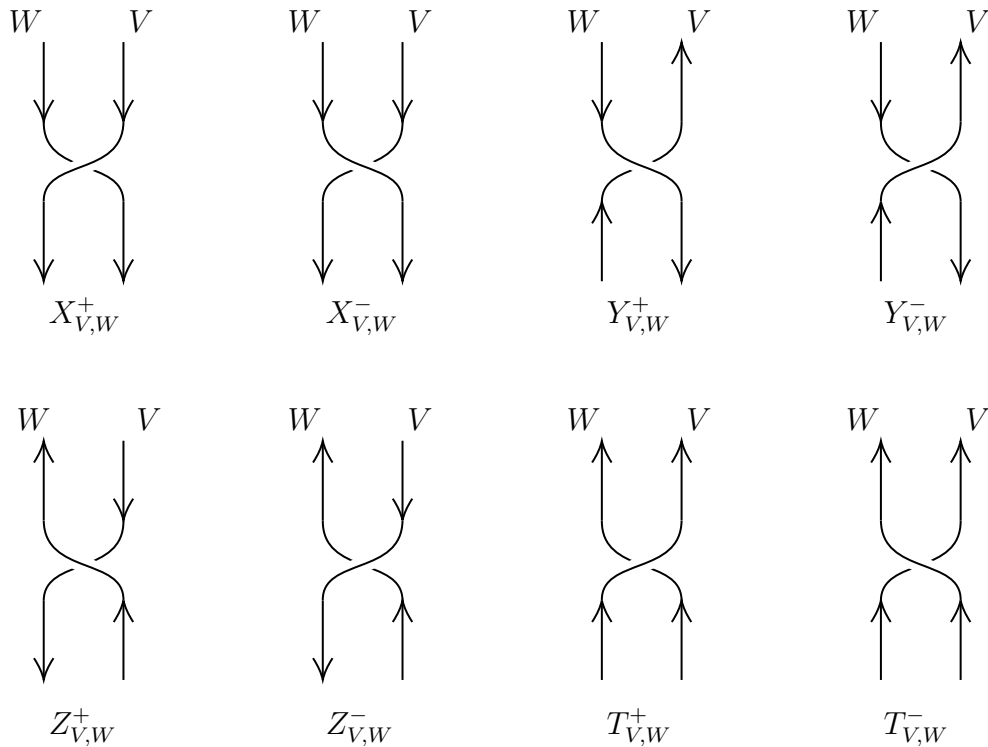
$$f: V_1^{\varepsilon_1} \otimes \cdots \otimes V_m^{\varepsilon_m} \rightarrow W_1^{\nu_1} \otimes \cdots \otimes W_n^{\nu_n}$$

where $V^{+1} = V$ and $V^{-1} = V^*$ for any object V of $\mathcal{V}$. Finally, we say that a ribbon graph is *v-colored* when it is colored and all its coupons are provided with colors as above.

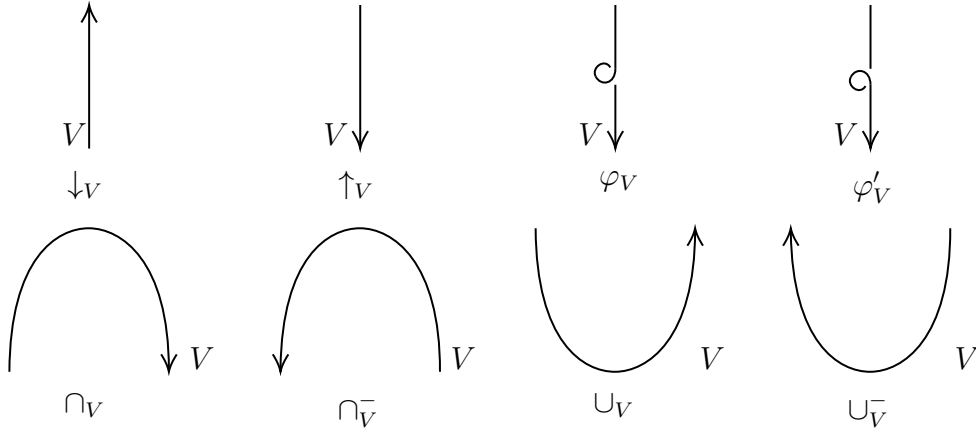
The *v*-colored graphs over $\mathcal{V}$ may be organized in a category denoted by $\mathbf{Rib}_{\mathcal{V}}$. The objects in $\mathbf{Rib}_{\mathcal{V}}$ are finite sequences $((V_1, \varepsilon_1), \dots, (V_m, \varepsilon_m))$ where $V_1, \dots, V_m$ are objects in $\mathcal{V}$ and $\varepsilon_1, \dots, \varepsilon_m \in \{-1, 1\}$. The empty sequence is also an object of $\mathbf{Rib}_{\mathcal{V}}$.

A morphism $\eta \rightarrow \eta'$ in $\mathbf{Rib}_{\mathcal{V}}$ is an isotopy type of a *v*-colored ribbon graph such that η (respectively η') is the sequence of colors and directions of those bands which hit the bottom (respectively the top) of the boundary intervals.

We may define composition of such morphisms in an analogous way we defined the composition of braids, by stacking a ribbon graph on top of the other, but of course, here, the number of bands that are being connects must agree, this follows from the definition of composition of morphisms, the target and source of the morphisms being composed must agree. To simplify notation, we may use the following notation on what follows:



Also, for the rest of distinguished morphisms we may have, we introduce the following notation:



A ribbon graph with no coupons is called a *ribbon tangle*. The ribbon tangles form a subcategory of $\mathbf{Rib}_{\mathcal{V}}$ which is also a strict monoidal category under the same tensor product.

Let's briefly recall that a covariant functor F between the categories $\mathcal{X}$ and $\mathcal{Y}$ associates to each object V of $\mathcal{X}$ an object $F(V)$ of $\mathcal{Y}$ and to each morphism $f: V \rightarrow W$ on $\mathcal{X}$ a morphism $F(f): F(V) \rightarrow F(W)$ on $\mathcal{Y}$, so that $F(\text{Id}_V) = \text{Id}_{F(V)}$. The covariance part of the definition is due to the property $F(f \circ g) = F(f) \circ F(g)$. If $\mathcal{X}$ and $\mathcal{Y}$ are monoidal categories, F is said to *preserve tensor product* if $F(1_{\mathcal{X}}) = 1_{\mathcal{Y}}$ and for any two objects or morphisms f, g of $\mathcal{X}$, we have $F(f \otimes g) = F(f) \otimes F(g)$.

Theorem 7.3.1. *Let $\mathcal{V}$ be a strict monoidal category with braiding c , twist θ and compatible duality $(*, b, d)$. There exists a unique covariant functor:*

$$F_{\mathcal{V}}: \mathbf{Rib}_{\mathcal{V}} \rightarrow \mathcal{V}$$

preserving tensor product and satisfying:

- (1) F transforms any object $(V, +1)$ into V and any object $(V, -1)$ into V^* ;
- (2) For any two objects V and W of $\mathcal{V}$, we have:

$$F(X_{V,W}^+) = c_{V,W}, F(\varphi_V) = \theta_V, F(\cup_V) = b_V, \text{ and } F(\cap_V) = d_V;$$

- (3) For any elementary v -colored ribbon graph Γ , we have $F(\Gamma) = f$ where f is the color of the only coupon of Γ .

Our objective in this chapter will be introduce the definition of the Topological Quantum Field Theories, and to do this, the proof of this theorem will not be necessary, although it can be seen in Turaev's book "*Quantum Invariants of Knots and 3-Manifolds*". What we need to know is only that such a functor exists. Also, this theorem establishes the connection we wanted between the ribbon categories and the study of ribbon graphs. This helps us to derive and understand properties on the ribbon categories, by analyzing the correspondent ribbon graphs.

Now, we introduce the remaining ingredients we need to finally give a definition for a TQFT.

7.4 Modular Tensor Categories

Definition 7.4.1. A category $\mathcal{V}$ is said to be a *Ab-category* if for every pair of V and W , the set $\text{Hom}(V, W)$ is an additive abelian group and the composition of morphisms is bilinear.

When talking about monoidal categories, we shall always assume that the tensor product is bilinear. Taking $\mathcal{V}$ a monoidal Ab-category, if we look at the composition of morphisms as a product in $\text{End}(\mathbb{1}) = \text{Hom}(\mathbb{1}, \mathbb{1})$, it gives this abelian group a structure of commutative ring with identity $\text{Id}_{\mathbb{1}}$. This ring is called the *ground ring* of $\mathcal{V}$, denoted by $\mathbb{K}_{\mathcal{V}}$ (or sometimes, when there is no confusion, only $\mathbb{K}$). Therefore, for any two objects V and W , the set $\text{Hom}(V, W)$ has the structure of a left $\mathbb{K}$ module when defining $kf = k \circ f$, for $k \in \mathbb{K}$ and $f \in \text{Hom}(V, W)$.

Definition 7.4.2. Let $\mathcal{V}$ be a ribbon Ab-category and let $\mathbb{K}_{\mathcal{V}} = \text{End}(\mathbb{1})$. An object V in $\mathcal{V}$ is called *simple* if $k \mapsto k \otimes \text{Id}_V$ defines a bijection $\mathbb{K} \rightarrow \text{End}(V)$.

Definition 7.4.3. Let $\{V_i\}_{i \in I}$ be a family of objects of a ribbon Ab-category $\mathcal{V}$. An object V of $\mathcal{V}$ is said to be *dominated* by $\{V_i\}_{i \in I}$ if there exists a finite set $\{V_{i(r)}\}_r$ and a family of morphisms $\{f_r: V_{i(r)} \rightarrow V, g_r: V \rightarrow V_{i(r)}\}_r$, such that:

$$\text{Id}_V = \sum_r f_r g_r.$$

Now, we can define modular categories, these will be useful to define invariants of closed oriented 3-manifolds and links embedded in such manifolds.

Definition 7.4.4. *Modular categories* are pairs $(\mathcal{V}, \{V_i\}_{i \in I})$ satisfying $\mathcal{V}$ is a ribbon Ab-Category and $\{V_i\}_{i \in I}$ is a finite family of simple objects satisfying:

- (M₁) (Normalization): There exists $0 \in I$ with $V_0 = \mathbb{1}$;
- (M₂) (Duality): For every $i \in I$ there exists $i^* \in I$ such that $V_{i^*} \simeq (V_i)^*$;
- (M₃) (Domination): For every object V in $\mathcal{V}$, V is dominated by $\{V_i\}_{i \in I}$;
- (M₄) (Non-Degeneracy): $S = [S_{i,j}]_{i,j \in I}$ is invertible over $\mathbb{K}$.

In the previous definition, the matrix S is constructed following

$$S_{i,j} = \text{tr}(c_{V_j, V_i} \circ c_{V_i, V_j}),$$

note that this matrix in fact belongs to $M_n(\mathbb{K})$, i.e., the set of all matrix $n \times n$ (n being the number of elements in I) with entries in $\mathbb{K} = \text{End}(\mathbb{1}, \mathbb{1})$. In fact, it is shown in [22] that this matrix is actually symmetric.

7.5 Invariants of 3-Manifolds

Consider L a framed link in S^3 and denote its components by $L_1, \dots, L_m$. Let U be closed tubular neighborhood of L in S^3 . We know U consists of m disjoint solid tori $U_1, \dots, U_m$ whose cores are the corresponding components of L .

Now we perform some surgery on S^3 in the following way: for each $n \in \{1, \dots, m\}$ we identify U_n with $S^1 \times B^2$ so that L_n is identified with $S^1 \times \{0\}$ and the framing of L_n is

identified with some normal vector field on this $S^1 \times \{0\}$. Now, for each copy of $S^1 \times B^2$, we glue a 2-handle $B^2 \times B^2$ to the B^4 bounded by S^3 (recall that a 3-sphere bounds a 4-ball), along the identifications:

$$\{U_n = S^1 \times B^2 = \partial B^2 \times B^2\}_{n=1}^m.$$

This gluing will obviously result in a compact connected 4-manifold, which we shall denote by W_L . The boundary of such a 4-manifold, ∂W_L , is a 3-manifold, composed by $S^3 - \text{Int}(U)$ and the m copies of $B^2 \times \partial B^2$ glued to $S^3 - \text{Int}(U)$ along the boundary. The obtained 3-manifold W_L is said to be obtained from surgery on S^3 along L .

Also, since W_L is a 4-manifold, it has a signature, coming from the signature of the intersection form $Q_{W_L}: H^2(W_L) \otimes H^2(W_L) \rightarrow \mathbb{Z}$ studied in the appendix. This signature may be denoted by $\sigma(W_L)$ or $\sigma(L)$, since it depends only on the link in which we are performing the surgery.

Consider now M a 3-manifold obtained by surgery on S^3 along a link L . Denote by $\text{col}(L)$ the set of all mappings from the set of components of L into I , then $|\text{col}(L)| = (|I|)^m$. In a way, we can understand $\text{col}(L)$ as the set of all possible colorings of the annuli that form L . For each $\lambda \in \text{col}(L)$, the pair (L, λ) determines a ribbon $(0, 0)$ -graph $\Gamma(L, \gamma)$ formed by m annuli. The cores are the components of L_i and their colors are $V_{\lambda(L_i)}$. Consider $F(\Gamma(L, \lambda)) \in \mathbb{K}$ where F is the functor defined earlier on Theorem 7.3.1, now set:

$$\{L\}_\lambda = \dim(\lambda)F(\Gamma(L, \lambda)) \in \mathbb{K}$$

where $\dim(\lambda) = \prod_{n=1}^m \dim(\lambda(L_n))$ and

$$\{L\} = \sum_{\lambda \in \text{col}(L)} \{L\}_\lambda \in \mathbb{K}.$$

Lemma 7.5.1. $\{L\}$ does not depend on the choice of orientation in L .

Proof. It suffices to prove that $\{L\}$ is preserved when the orientation of one component is changed. Call L' the link obtained from L when only one component, L_n , of L have its orientation reversed. For a certain coloring $\lambda \in \text{col}(L)$, denote by λ' the coloring of L' which coincides with λ on all components of $L - L'$, and on the component that had the orientation reversed it equals $(\lambda(L_n))^* \in I$.

We know that $\dim(V_i) = \dim(V_{i*})$ for every object V_i , hence, $\{L'\}_{\lambda'} = \{L\}_\lambda$, therefore the summing over all these terms result in the same value, thus $\{L\} = \{L'\}$. $\square$

In order to define the invariant on 3-manifolds, we shall introduce yet two other elements. The first will be denoted as $\mathcal{D}$ and called the *rank* of $\mathcal{V}$. We define $\mathcal{D}$ as an element of the ground ring $\mathbb{K}_{\mathcal{V}}$ given by:

$$\mathcal{D}^2 = \sum_{i \in I} (\dim(i))^2.$$

(we denote $\dim(i) = \dim(V_i)$). In general, $\mathcal{D}$ may not exist, and when it exists, it may not be unique.

The second element, is denoted by $\Delta_{\mathcal{V}}$ given by:

$$\Delta_{\mathcal{V}} = \sum_{i \in I} v_i^{-1} (\dim(i))^2 \in \mathbb{K}_{\mathcal{V}}.$$

Here, since V_i is a simple object, the twist on $\mathcal{V}$ acts on V_i as multiplication by a certain $v_i \in \mathbb{K}$, and as we've seen the twist acts as via isomorphisms, hence the v_i 's have inverses v_i^{-1} in $\mathbb{K}$.

Now we can define an invariant of 3-manifolds using all these notions we have defined throughout this chapter, we define:

$$\tau(M) = \tau_{\mathcal{V}, \mathcal{D}}(M) = \Delta^{\sigma(L)} \mathcal{D}^{-\sigma(L)-m-1} \{L\} \in \mathbb{K}.$$

Theorem 7.5.1. $\tau(M)$ is a topological invariant of M .

We may even extend this invariant to invariants of ribbon graphs embedded in 3-manifolds. The definition of ribbon $(0, 0)$ -graph extends to arbitrary 3-manifolds. Let M be a closed connected oriented 3-manifold and Ω a v -colored ribbon graph over $\mathcal{V}$ in M . Let's think of M as being the result of surgery on S^3 along a link L with components $L_1, \dots, L_m$. Now we fix an orientation for L . Let U be the tubular neighborhood of L , then by using isotopy, we can deform Ω inside M in such a way to avoid U , thus we may assum, in general, that $\Omega \subset S^3 - U$.

For any $\lambda \in \text{col}(L)$, we form a v -colored $(0, 0)$ -graph $\Gamma(L, \lambda) \subset U$, the union $\Gamma(L, \lambda) \cup \Omega$ is again a v -colored ribbon $(0, 0)$ -graph in S^3 . Let:

$$\{L, \Omega\} = \sum_{\lambda \in \text{col}(L)} \dim(\lambda) F(\Gamma(L, \lambda) \cup \Omega) \in \mathbb{K}.$$

Lemma 7.5.2. $\{L, \Omega\}$ does not depend on the choice of orientation in L .

The proof of such a lemma is just a modification of the proof of Section 7.5. And the most general result, that gives us an invariant of ribbon graphs inside 3-manifolds, is concerning the following element:

$$\tau(M, \Omega) = \tau_{\mathcal{V}, \mathcal{D}}(M, \Omega) = \Delta^{\sigma(L)} \mathcal{D}^{-\sigma(L)-m-1} \{L, \Omega\}.$$

Theorem 7.5.2. $\tau(M, \Omega)$ is a topological invariant of the pair (M, Ω) .

These invariants will appear in the future when we talk about three dimensional TQFT's, they will provide us with a almost ready TQFT for general 3-manifolds.

The invariant τ extends to v -colored ribbon graphs in any non-connected closed oriented 3-manifold M by the formula:

$$\tau(M, \Omega) = \prod_r \tau(M_r, \Omega_r)$$

where M_r are the connected components of M and Ω_r the part of Ω lying inside M_r .

7.6 The definition of a TQFT

Space-Structures

Let's start by introducing a generalized form of a cobordism, to this we first generalize the notion of orientation and define the space structures on manifolds. To have a first sense on what a space structure is, let's consider n a positive integer, X an n -dimensional topological manifold and $\text{Ori}_n(X)$ the set of all orientations in X . If we consider any homeomorphism $f: X \rightarrow Y$, it is clear that f induces a bijection $\text{Ori}_n(X) \rightarrow \text{Ori}_n(Y)$,

since f being a homeomorphism, Y have all the same possible orientations as X , as well as the same dimension let's, not coincidentally denote such a bijection as $\text{Ori}_n(f)$. Then $\text{Ori}_n(\text{Id}_X) = \text{Id}_{\text{Ori}_n(X)}$ and $\text{Ori}_n(f \circ g) = \text{Ori}_n(f) \circ \text{Ori}_n(g)$, or, in words, the formulas $X \mapsto \text{Ori}_n(X)$ and $f \mapsto \text{Ori}_n(f)$ define a covariant functor from the category of topological manifolds and homeomorphisms to the category of sets and bijections.

This is a specific case of a space structure, which we shall now formally define. A *space structure* is a covariant functor from the category of topological spaces and homeomorphisms to the category of sets and bijections, and satisfying that its value on the empty space is a singleton. Let $\mathfrak{A}$ be such a functor, the definition of space structure says that if X and Y are topological spaces and $f: X \rightarrow Y$ a homeomorphism, then $\mathfrak{A}(X)$ is a set, $\mathfrak{A}(f): \mathfrak{A}(X) \rightarrow \mathfrak{A}(Y)$ is a bijection and for every homeomorphism $g: Y \rightarrow Z$, we have $\mathfrak{A}(f \circ g) = \mathfrak{A}(f) \circ \mathfrak{A}(g)$. Also, we note that $\mathfrak{A}(\text{Id}_X) = \text{Id}_{\mathfrak{A}(X)}$.

Elements of the set $\mathfrak{A}(X)$ are called $\mathfrak{A}$ -structures on X , and any pair (X, α) , with $\alpha \in \mathfrak{A}(X)$ is called a $\mathfrak{A}(X)$ -space. A $\mathfrak{A}$ -homeomorphism of a $\mathfrak{A}$ -space (X, α) onto an $\mathfrak{A}$ -space (X', α') is a homeomorphism $f: X \rightarrow X'$, such that $\mathfrak{A}(f)(\alpha) = \alpha'$.

We will also need that the space structure be compatible with disjoint union of topological spaces, that is to say that for every topological spaces X and Y , we have a canonical map $\mathfrak{A}(X) \times \mathfrak{A}(Y) \rightarrow \mathfrak{A}(X \amalg Y)$ satisfying the following properties, given in a diagrammatically way:

- (1) The diagram commutes:

$$\begin{array}{ccc} \mathfrak{A}(X) \times \mathfrak{A}(Y) & \longrightarrow & \mathfrak{A}(X \amalg Y) \\ \downarrow \text{Perm} & & \downarrow = \\ \mathfrak{A}(Y) \times \mathfrak{A}(X) & \longrightarrow & \mathfrak{A}(Y \amalg X) \end{array}$$

- (2) For every $f: X \rightarrow X'$ and $g: Y \rightarrow Y'$, we have commutative:

$$\begin{array}{ccc} \mathfrak{A}(X) \times \mathfrak{A}(Y) & \longrightarrow & \mathfrak{A}(X \amalg Y) \\ \downarrow \mathfrak{A}(f) \times \mathfrak{A}(g) & & \downarrow \mathfrak{A}(f \amalg g) \\ \mathfrak{A}(X') \times \mathfrak{A}(Y') & \longrightarrow & \mathfrak{A}(X' \amalg Y') \end{array}$$

- (3) The diagram commutes:

$$\begin{array}{ccc} \mathfrak{A}(X) \times \mathfrak{A}(Y) \times \mathfrak{A}(Z) & \longrightarrow & \mathfrak{A}(X \amalg Y) \times \mathfrak{A}(Z) \\ \downarrow & & \downarrow \\ \mathfrak{A}(X) \times \mathfrak{A}(Y \amalg Z) & \longrightarrow & \mathfrak{A}(X \amalg Y \amalg Z) \end{array}$$

- (4) The canonical mapping $\mathfrak{A}(\emptyset) \times \mathfrak{A}(Y) \rightarrow \mathfrak{A}(\emptyset \amalg Y) = \mathfrak{A}(Y)$ is induced by the identity $\text{Id}_{\mathfrak{A}(Y)}: \mathfrak{A}(Y) \rightarrow \mathfrak{A}(Y)$.

Recall that an involution on a set $\mathfrak{A}(X)$ is a map $\iota: \mathfrak{A}(X) \rightarrow \mathfrak{A}(X)$ such that $\iota \circ \iota = \text{Id}_{\mathfrak{A}(X)}$. Thus we say that a space structure $\mathfrak{A}$ is involutive when for every topological space X , the set $\mathfrak{A}(X)$ is equipped with an involution, and such that:

- (i) The canonical map $m_{X,Y}: \mathfrak{A}(X) \times \mathfrak{A}(Y) \rightarrow \mathfrak{A}(X \amalg Y)$ is equivariant with respect to the involutions on X , Y and $X \amalg Y$, i.e., $m_{X,Y} \circ (\iota_X \times \iota_Y) = \iota_{X \amalg Y} \circ m_{X,Y}$;

- (ii) $\mathfrak{A}(f): \mathfrak{A}(X) \rightarrow \mathfrak{A}(Y)$ is equivariant with respect to the involutions of X and Y , i.e., $\mathfrak{A}(f) \circ \iota_X = \iota_Y \circ \mathfrak{A}(f)$.

Also, if we denote by $-\alpha$ the image of a space structure α in X via the involution ι_X , then we denote the $\mathfrak{A}(X)$ -space $(X, -\alpha)$ as $-X$ and say that $-\alpha$ is the *opposite* space structure of α . For a $\mathfrak{A}$ -homeomorphism $f: X \rightarrow Y$, $-f$ is the same homeomorphism f , but between the opposite $\mathfrak{A}$ -spaces, i.e., $-f: -X \rightarrow -Y$.

Modular Functors

Consider fixed a space structure $\mathfrak{A}$ compatible with disjoint union and consider K a commutative ring with identity. We define a *modular functor* $\mathcal{T}$ with ground ring K based on the space structure $\mathfrak{A}$ as a functor that assigns to every $\mathfrak{A}$ -space X a projective K -module of finite type $\mathcal{T}(X)$ and to every $\mathfrak{A}$ -homeomorphism $f: X \rightarrow Y$ a K -isomorphism $f_\#: \mathcal{T}(X) \rightarrow \mathcal{T}(Y)$ satisfying the following properties:

- (T₁) For any $\mathfrak{A}$ -homeomorphisms $f: X \rightarrow Y$, $g: Y \rightarrow Z$, we have $(g \circ f)_\# = g_\# \circ f_\#$;
- (T₂) For disjoint $\mathfrak{A}$ -spaces X and Y there is an identification isomorphism $\mathcal{T}(X \amalg Y) = \mathcal{T}(X) \otimes_K \mathcal{T}(Y)$ satisfying the following three conditions:
- (i) (Commutativity): the commutativity of the diagram

$$\begin{array}{ccc} \mathcal{T}(X \amalg Y) & \xrightarrow{=} & \mathcal{T}(X) \otimes_K \mathcal{T}(Y) \\ \downarrow = & & \downarrow \text{Perm.} \\ \mathcal{T}(Y \amalg X) & \xrightarrow{=} & \mathcal{T}(Y) \otimes_K \mathcal{T}(X) \end{array}$$

- (ii) (Associativity): $(\mathcal{T}(X) \otimes_K \mathcal{T}(Y)) \otimes \mathcal{T}(Z) = (X \amalg Y) \otimes_K \mathcal{T}(Z) = \mathcal{T}(X \amalg Y \amalg Z) = \mathcal{T}(X) \otimes_K (\mathcal{T}Y \amalg Z) = \mathcal{T}(X) \otimes_K (\mathcal{T}(Y) \otimes_K \mathcal{T}(Z))$ is the identification $(x \otimes y) \otimes z = x \otimes (y \otimes z)$.
- (iii) (Naturality): For every $f: X \rightarrow X'$ and $g: Y \rightarrow Y'$, the diagram commutes:

$$\begin{array}{ccc} \mathcal{T}(X \amalg Y) & \xrightarrow{(f \amalg g)_\#} & \mathcal{T}(X' \amalg Y') \\ \downarrow = & & \downarrow = \\ \mathcal{T}(X) \otimes_K \mathcal{T}(Y) & \xrightarrow{f_\# \otimes g_\#} & \mathcal{T}(X') \otimes_K \mathcal{T}(Y') \end{array}$$

- (T₃) $\mathcal{T}(\emptyset) = K$ which induces for any $\mathfrak{A}$ -space Y , the identification $\mathcal{T}(Y) = \mathcal{T}(\emptyset \amalg Y) = \mathcal{T}(\emptyset) \otimes_K \mathcal{T}(Y)$.

The K -module $\mathcal{T}(X)$ is called the *module of states* of X and its elements are, therefore called the $\mathcal{T}$ -states on X .

As an almost trivial example, we may consider the constant functor $\mathcal{T}(X) = K$ for every topological space X and $\mathcal{T}(f) = \text{Id}_X$ for every homeomorphism $f: X \rightarrow Y$. Then $\mathcal{T}$ is clearly a modular functor since it satisfies all the three properties established before.

Suppose the modular functor $\mathcal{T}$ is based on an involutive structure space $\mathfrak{A}$, then if for every $\mathfrak{A}$ -space X we can find a non-degenerate pairing $d_X: \mathcal{T}(X) \otimes_K \mathcal{T}(-X) \rightarrow K$ and the system of pairings $\{d_X\}_X$ is natural with respect to $\mathfrak{A}$ -homeomorphisms (i.e., $d_X = d_Y \circ (f_\# \otimes (-f)_\#)$), multiplicative with respect to disjoint unions (i.e., $d_{X \amalg Y} = d_X \cdot d_Y$) and symmetric in the sense that $d_{-X} = d_X \circ \text{Perm}_{\mathcal{T}(-X) \otimes_K \mathcal{T}(X)}$ for any $\mathfrak{A}$ -space, then we say that $\mathcal{T}$ is *self dual*.

7.6.1 Cobordisms of $\mathfrak{A}$ -spaces

In the realm of manifolds we will see that $(n+1)$ -TQFT's associates modules to closed n -manifolds and homomorphisms to cobordisms between n -manifolds.

We will say that an $(n+1)$ -cobordism is a compact $(n+1)$ -manifold whose boundary is split as a disjoint union of two closed manifolds called the bases of the cobordism:

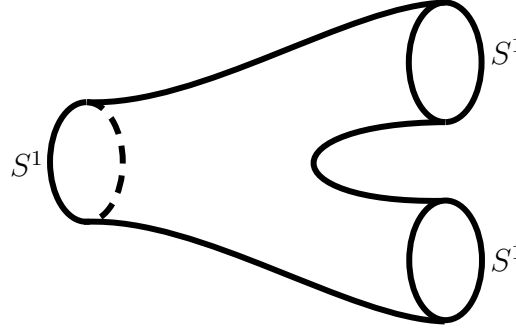


Figure 7.1: A cobordism between S^1 and $S^1 \amalg S^1$.

But before we restrict ourselves to this case, we shall define a more general notion of a cobordism using what we defined earlier on space structures. Let $\mathfrak{A}$ and $\mathfrak{B}$ be space structures compatible with disjoint union. We suppose that $\mathfrak{A}$ is involutive, but do not require the same thing for $\mathfrak{B}$. Suppose now that any $\mathfrak{B}$ -space M contains a topological subspace with an $\mathfrak{A}$ -structure. We will call this subspace the *boundary* of M and denote it as ∂M .

We may assume that the boundary is natural with respect to $\mathfrak{B}$ -homeomorphisms, i.e., $\mathfrak{B}$ -homeomorphism $M_1 \rightarrow M_2$ that restricts to an $\mathfrak{A}$ -homeomorphism $\partial M_1 \rightarrow \partial M_2$. We also assume that the boundary commutes with disjoint union, i.e., $\partial(M_1 \amalg M_2) = \partial M_1 \amalg \partial M_2$.

Definition 7.6.1. The pair $(\mathfrak{B}, \mathfrak{A})$ forms a *cobordism theory* if it satisfies:

- (CB₁) The $\mathfrak{B}$ -spaces are subject to gluing as follows: Let M be a $\mathfrak{B}$ -space whose boundary is $X \amalg Y \amalg Z$, $\mathfrak{A}$ -spaces such that X is an $\mathfrak{A}$ -homeomorphic to $-Y$. Let M' be the topological space obtained by gluing X to Y along the homeomorphism $f: X \rightarrow -Y$. The $\mathfrak{B}$ -structure on M gives a $\mathfrak{B}$ -structure on M' , such that $\partial M' = Z$. The $\mathfrak{B}'$ -structures on M' corresponding in this way to $f: X \rightarrow -Y$ and $(-f)^{-1}: Y \rightarrow -X$ coincide;
- (CB₂) The gluing of $\mathfrak{B}$ -spaces are natural with respect to $\mathfrak{B}$ -homeomorphisms and commute with the disjoint union. The gluing along disjoint $\mathfrak{A}$ -subspaces of the boundary coincides with the gluing along their union;
- (CB₃) Let X and X' be two copies of the $\mathfrak{A}$ -space. Gluing $\mathfrak{B}$ -spaces $X \times [0, 1]$ and $X' \times [0, 1]$ along the identity $X \times \{1\} \rightarrow X' \times \{0\}$ given by $(x, 1) \mapsto (x, 0)$ gives us a $\mathfrak{B}$ -space homeomorphic to the cylinder $X \times [0, 1]$ via a $\mathfrak{B}$ -homeomorphism which is the identity on the bases.

As we said earlier, taking $\mathfrak{B}$ the structure of orientations on compact smooth $(n+1)$ -manifolds possibly with boundary, and $\mathfrak{A}$ the structure of orientations of n -manifolds closed and smooth, ∂ the usual boundary, we recover the cobordism theory that we have mentioned, which is the more standard cobordism theory used.

7.6.2 Definition of Topological Quantum Field Theory

We have finally defined all the concepts we need to define in a quite general fashion the concept of a Topological Quantum Field Theory, or TQFT for short. Consider first $\mathfrak{A}$ an involutive space-structure and $\mathfrak{B}$ a space structure. Assume also that the pair $(\mathfrak{B}, \mathfrak{A})$ forms a cobordism theory. Then a $(\mathfrak{B}, \mathfrak{A})$ -cobordism is a triple (M, X, Y) where M is a $\mathfrak{B}$ -space, X and Y are $\mathfrak{A}$ -spaces and $\partial M = (-X) \amalg Y$.

The $\mathfrak{A}$ -spaces X and Y are called the *bottom* and *top* bases of this cobordism and denoted by $\partial_- M$ and $\partial_+ M$ respectively. We also say that M is a cobordism between X and Y . If we go back to Figure 7.1, we have that M is the 2-manifold sometimes called *the pair of pants* (due to the fact that it looks like a pair of pants) and its bases are $\partial_- M = S^1$ and $\partial_+ M = S^1 \amalg S^1$, of course, $\partial M = S^1 \amalg S^1 \amalg S^1$.

We may also glue together cobordisms M_1 and M_2 along any $\mathfrak{A}$ -homeomorphism $\partial_+ M_1 \rightarrow \partial_- M_2$. In Figure 7.1 the picture of the pair of pants was represented with its bottom base on the left and its top base on the right, if we want to make it consistent with the nomenclature given, we should also draw the same picture but with the bottom bases on the bottom and the top bases on the top, the figure would, then, look like an upside down pair of pants.

The gluing between two cobordisms would pictorially look like the stacking of a cobordism on top of the other, in the next picture, we have an example of the gluing of the pair of pants $(M, S^1 \amalg S^1, S^1)$ and a disk, which can be seen as the cobordism $(D^2, S^1, \emptyset)$:

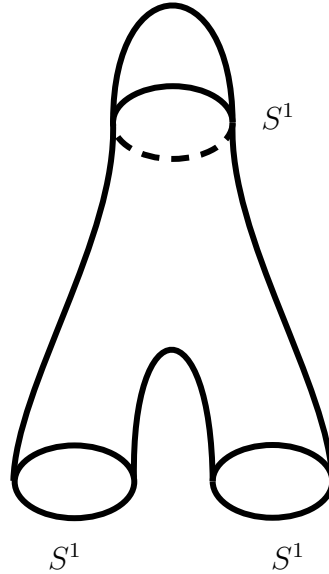


Figure 7.2: The gluing of two cobordisms.

A *Topological Quantum Field Theory* based on the cobordism theory $(\mathfrak{B}, \mathfrak{A})$ consists of a modular functor $\mathcal{T}$ with ground ring K based on $\mathfrak{A}$ (recall that a modular functor needs a space structure to be defined on) and a map τ which assigns to every $(\mathfrak{B}, \mathfrak{A})$ -cobordism (M, X, Y) a K -homomorphism

$$\tau(M) = \tau(M, X, Y): \mathcal{T}(X) \rightarrow \mathcal{T}(Y)$$

which satisfies:

(TQFT₁): (Naturality) If M_1 and M_2 are $(\mathfrak{B}, \mathfrak{A})$ -cobordisms and $f: M_1 \rightarrow M_2$ is a $\mathfrak{B}$ -homeomorphism preserving the bases, then the diagram commutes:

$$\begin{array}{ccc} \mathcal{T}(\partial_- M_1) & \xrightarrow{\tau(M_1)} & \mathcal{T}(\partial_+ M_1) \\ (f|_{\partial_- M_1})_{\#} \downarrow & & \downarrow (f|_{\partial_+ M_2})_{\#} \\ \mathcal{T}(\partial_- M_2) & \xrightarrow{\tau(M_2)} & \mathcal{T}(\partial_+ M_2) \end{array}$$

(TQFT₂): (Multiplicativity) If a cobordism M is the disjoint union of cobordisms M_1 and M_2 then, under identifications, we have $\tau(M) = \tau(M_1) \otimes \tau(M_2)$;

(TQFT₃): (Functoriality) If a $(\mathfrak{B}, \mathfrak{A})$ -cobordism M is obtained from disjoint union of two $(\mathfrak{B}, \mathfrak{A})$ -cobordisms M_1 and M_2 by gluing along $\mathfrak{A}$ -homeomorphism $f: \partial_+ M_1 \rightarrow \partial_- M_2$ then for some invertible $k \in K$:

$$\tau(M) = k\tau(M_2) \circ f_{\#} \circ \tau(M_1);$$

(TQFT₄): (Normalization) For any $\mathfrak{A}$ -space X , we have

$$\tau(X \times [0, 1], X \times \{0\}, X \times \{1\}) = \text{Id}_{\mathcal{T}(X)}.$$

The homomorphism $\tau(M, X, Y): \mathcal{T}(X) \rightarrow \mathcal{T}(Y)$ is called *operator invariant* of the cobordism M . The factor k that appears on the functoriality axiom is called *anomaly* and when we may always take $k = 1$ we call the TQFT an *anomaly-free* TQFT. Note also, that the anomaly of a TQFT need not be unique, for example, whenever $\tau(M_1) = \tau(M_2) = 0$, we can take any invertible value of k , so that (TQFT₃) is satisfied.

Every $\mathfrak{B}$ -space M gives rise to a cobordism $(M, \emptyset, \partial M)$, and the corresponding K -homomorphism given by the operator invariant $\tau(M): \mathcal{T}(\emptyset) \rightarrow \mathcal{T}(\partial M)$ is completely determined by its value on the identity $1_K \in K = \mathcal{T}(\emptyset)$, we shall denote this value as $\tau(M)$ and consequently $\tau(M) \in \mathcal{T}(\partial M)$. The following picture shows in a more specific context, namely that when the space structures taken are the spaces of orientation, and the cobordisms are between 1-dimensional manifolds, how we obtain, via homeomorphism, the cobordism $(M, \emptyset, \partial M)$ from $(M, \partial_- M, \partial_+ M)$:

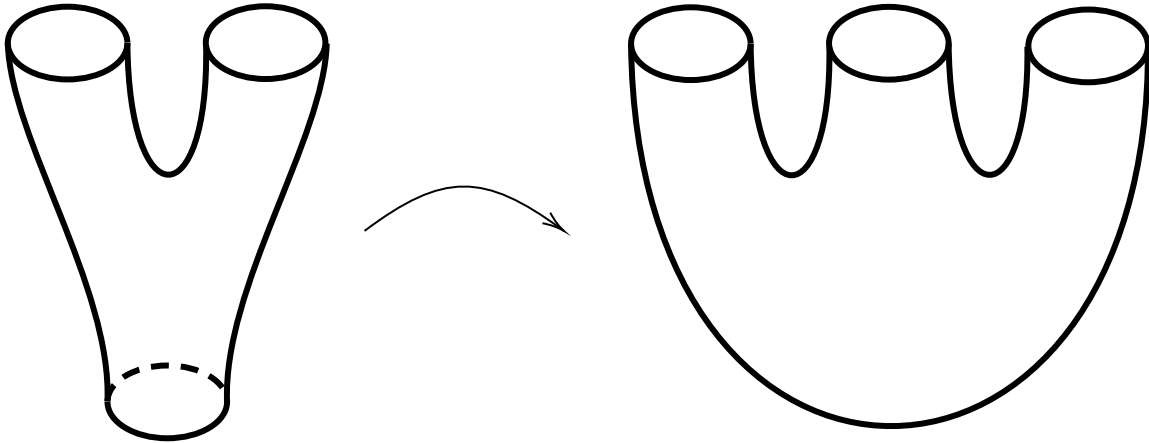


Figure 7.3: How to obtain $(M, \emptyset, S^1 \amalg S^1 \amalg S^1)$ from $(M, S^1, S^1 \amalg S^1)$

We shall say that a $\mathfrak{B}$ -space is *closed* when $\partial M = \emptyset$. For a closed $\mathfrak{B}$ -space M , we have $\tau(M) \in \mathcal{T}(\emptyset) = K$, so that $\tau(M)$ is a K -valued $\mathfrak{B}$ -homeomorphism invariant of M .

We may also define tensor products of TQFT's in the following way, consider $(\mathcal{T}_1, \tau_1)$ and $(\mathcal{T}_2, \tau_2)$ two TQFT's, their tensor product is defined by $(\mathcal{T}_1 \otimes \mathcal{T}_2)(X) = \mathcal{T}_1(X) \otimes \mathcal{T}_2(X)$ and $(\tau_1 \otimes \tau_2)(M) = \tau_1(M) \otimes \tau_2(M)$, thus the new TQFT $(\mathcal{T}_1 \otimes \mathcal{T}_2, \tau_1 \otimes \tau_2)$ is a new TQFT over the cobordism $(\mathfrak{B}, \mathfrak{A})$, since the axioms of TQFT follow directly from the fact that $(\mathcal{T}_1, \tau_1)$ and $(\mathcal{T}_2, \tau_2)$ satisfy them.

Referências

- [1] J. S . Birman. Braids, link polynomials and a new algebra. *transactions of the american mathematical society*, 17:249–273, 1989.
- [2] JOAN S. BIRMAN and James Cannon. *Braids, Links, and Mapping Class Groups. (AM-82)*. Princeton University Press, 1974.
- [3] G.E. Bredon, J.H. Ewing, F.W. Gehring, and P.R. Halmos. *Topology and Geometry*. Graduate Texts in Mathematics. Springer, 1993.
- [4] Brian J. Sanderson Colin P. Rourke. *Introduction to piecewise linear topology*. Springer Berlin, Heidelber, The University of Warwick, England, 1 edition, 1982.
- [5] Anthony Conway, Stefan Friedl, and Gerrit Herrmann. Linking forms revisited, 2017.
- [6] Louis Funar. *A brief introduction to knot invariant*. Pre Print, Grenoble, version 2.0 edition, 2014.
- [7] R.E. Gompf and A. Stipsicz. *4-Manifolds and Kirby Calculus*. Graduate studies in mathematics. American Mathematical Society, 1999.
- [8] Allen Hatcher. *Algebraic topology*. Cambridge University Press, Cambridge, 2002.
- [9] J. Hempel. *3-Manifolds*. AMS Chelsea Publishing Series. AMS Chelsea Pub., American Mathematical Society, 2004.
- [10] S.T. Hu. *Homotopy Theory*. Homotopy Theory. Department of Mathematics, Tulane University, 1950.
- [11] J. HUMPHREYS. *Introduction to Lie Algebras and Representation Theory*. Graduate Texts in Mathematics. Springer New York, 1994.
- [12] V. F. R. Jones. Hecke algebra representations of braid groups and link polynomials. *Annals of Mathematics*, 126:335–388, 1987.
- [13] Akio Kawauchi and Sadayoshi Kojima. Algebraic classification of linking pairings on 3-manifolds. *Mathematische Annalen*, 253:29–42, 1980.
- [14] GREG KUPERBERG. The quantum g_2 link invariant. *International Journal of Mathematics*, 05(01):61–85, 1994.
- [15] W. B. Raymond Lickorish. *An Introduction to Knot Theory*. Springer New York, NY, Cambridge, England, 1 edition, 1997.

- [16] Edwin E. Moise. *Geometric Topology in Dimensions 2 and 3*. Springer New York, NY, Department of Mathematics, Queens College, CUNY, Flushing, USA, 1 edition, 1977.
- [17] J.R. Munkres. *Elements Of Algebraic Topology*. CRC Press, 2018.
- [18] J. Rotman. *An Introduction to Algebraic Topology*. Graduate Texts in Mathematics. Springer New York, 1998.
- [19] N. Saveliev. *Lectures on the Topology of 3-Manifolds: An Introduction to the Casson Invariant*. De Gruyter Textbook. De Gruyter, 2012.
- [20] M. Spivak. *Calculus On Manifolds: A Modern Approach To Classical Theorems Of Advanced Calculus*. Avalon Publishing, 1965.
- [21] Dennis Sullivan. On the intersection ring of compact three manifolds. *Topology*, 14(3):275–277, 1975.
- [22] V.G. Turaev. *Quantum Invariants of Knots and 3-Manifolds*. De Gruyter Studies in Mathematics. De Gruyter, 2016.
- [23] Vladimir Turaev. The yang-baxter equation and invariants of links. *Inventiones Mathematicae*, 92:527–553, 10 1988.
- [24] Vladimir Turaev. Cohomology rings, linking forms and invariants of spin structures of three-dimensional manifolds. *Mathematics of the USSR-Sbornik*, 48:65, 10 2007.
- [25] James W. Vick. *Homology Theory*. Springer New York, NY, Austin, USA, 2 edition, 1994.
- [26] G.W. Whitehead. *Elements of Homotopy Theory*. Graduate Texts in Mathematics. Springer New York, 1978.

A A Review on Homology and Cohomology

A.1 General Review

Homology

In this chapter we review some basic definitions and results on homology and cohomology theory. We emphasize that since these results were done in previous work, we shall not prove all of them, only those which gives us some intuition on how to work with given definitions.

We start by revisiting the definition of homology groups. Here, we denote by σ_p the **standard p -simplex**, i.e., the convex hull of the points $e_0 = (1, 0, \dots, 0)$, $e_1 = (0, 1, \dots, 0)$, $\dots$, $e_p = (0, 0, \dots, 1)$. We define a **singular p -simplex** as a function $\varphi: \sigma_p \rightarrow X$ where X is a topological space. If φ is a singular p -simplex, we define the **i -th face of φ** as the singular $(p-1)$ -simplex $\partial_i \varphi$ given by:

$$\partial_i \varphi(t_0, \dots, t_{p-1}) = \varphi(t_0, \dots, t_{i-1}, 0, t_i, \dots, t_{p-1})$$

in other words, $\partial_i \varphi$ is the singular $(p-1)$ -simplex obtained from φ by omitting its i -th coordinate. Recall that σ_p is the set of all n -tuples $(t_0, \dots, t_p) \in \mathbb{R}^{p+1}$ such that $\sum_{i=1}^p t_i = 1$ and $t_i \geq 0$ for all i .

Let $f: X \rightarrow Y$ be a continuous function from the topological space X to the topological space Y and φ a singular p -simplex in X , then we can define an **induced** singular p -simplex in Y by $f_{\#}(\varphi) = f \circ \varphi$. The function $f_{\#}$ satisfies

$$(g \circ f)_{\#} = g_{\#} \circ f_{\#} \text{ and } \text{Id}_{\#} = \text{Id}.$$

For a topological space X we can define the group $S_n(X)$ as the free abelian group such that its basis is the set of all singular p -simplexes in X . An element of X is therefore a linear combination

$$\sum_{\varphi} n_{\varphi} \varphi$$

with $n_{\varphi} \in \mathbb{Z}$, these elements are called **singular p -chains**.

On $S_n(X)$ we can define something called **boundary operator** which will be denoted as ∂ , and it takes an element of $S_n(X)$ to an element in S_{n-1} it does this by means of the faces of the simplexes forming an element in $S_n(X)$: first we define $\partial_i: S_n(X) \rightarrow S_{n-1}(X)$ given by:

$$\partial_i \left(\sum_{\varphi} n_{\varphi} \varphi \right) = \sum_{\varphi} n_{\varphi} \partial_i \varphi$$

now, we define $\partial: S_n(X) \rightarrow S_{n-1}(X)$ by setting:

$$\partial = \sum_{i=1}^n (-1)^i \partial_i.$$

With simple computations, we can see that $\partial \circ \partial = 0$.

Now, let's name some important elements in $S_n(X)$. If $c \in S_n(X)$ is an element in the kernel of ∂ , i.e., $\partial c = 0$, then it is called an **n -cycle**. If d is an element of $S_n(X)$ that is the image of some element in $S_{n+1}(X)$ via ∂ , then d is called an **n -boundary**, hence there must exist $e \in S_{n+1}(X)$ such that $d = \partial e$. Now it is clear that every n -boundary is an n -cycle since $\partial \circ \partial = 0$, more than that, if we define $Z_n(X)$ as the kernel of $\partial: S_n(X) \rightarrow S_{n-1}(X)$ and $B_n(X)$ as the image of $\partial: S_{n+1}(X) \rightarrow S_n(X)$, then these sets are obviously groups and $B_n(X)$ is a normal subgroup of $Z_n(X)$, thus we can finally define the **Homology Group** of X as being the quotient:

$$H_n(X) = \frac{Z_n(X)}{B_n(X)}$$

From the definition of $S_n(X)$ and ∂ we may see $(S_*(X), \partial)$ as a chain complex, thus $H_*(X)$ is a graded group. Also if we take $(S_*(X'), \partial')$ another chain complex and $\phi: S_*(X) \rightarrow S_*(X')$ a chain homomorphism, then we say that ϕ is a chain map satisfying $\partial' \circ \phi = \phi_{n-1} \partial$. Of course, a chain map like this induces a homomorphism on the homology groups $\phi_*: H_*(X) \rightarrow H_*(X')$.

Lemma A.1.1. *If X is a topological space consisting of one point then $H_0(X) \simeq \mathbb{Z}$ and $H_n(X) \simeq 0$ for all $n > 0$.*

Proof. For all $n \geq 0$ a singular n -simplex is a constant function taking each point on σ_n to $p \in X$, thus $S_n(X) = \{k\varphi_n; k \in \mathbb{Z}\}$ with φ being such a constant function. Considering the chain complex $(S_*(X), \partial)$, since each $S_n(X)$ is generated by φ_n , we get:

$$\partial \varphi_n = \sum_{i=0}^n (-1)^i \partial_i \varphi_n = \sum_{i=1}^n (-1)^i \varphi_{n-1}$$

thus, $\partial \varphi_{2n-1} = 0$ and $\partial \varphi_{2n} = \varphi_{2n-1}$. Now computing $Z_n(X)$ and $B_n(X)$ we get:

$$Z_n(X) = \begin{cases} S_n(X), & \text{if } n \text{ is odd or } 0 \\ 0, & \text{if } n \text{ is even} \end{cases}$$

$$B_n(X) = \begin{cases} S_n(X), & \text{if } n \text{ is odd} \\ 0, & \text{if } n \text{ is even} \end{cases}$$

and hence $\frac{Z_0(X)}{B_0(X)} \simeq S_0(X) \simeq \mathbb{Z}$ and $\frac{Z_n(X)}{B_n(X)} \simeq \frac{S_n(X)}{S_n(X)} = 0$. □

Theorem A.1.1. *If $f: X \rightarrow Y$ is a homeomorphism, then $f_*: H_p(X) \rightarrow H_p(Y)$ is an isomorphism.*

Recall that when a sequence of homomorphism and groups $C \xrightarrow{f} D \xrightarrow{g} E$ satisfies $\text{im}(f) = \ker(g)$ such a sequence is called **exact**. If the sequence $0 \rightarrow C \xrightarrow{f} D \xrightarrow{g} E \rightarrow 0$ is exact, then we call it **short exact**.

Consider $C = \{C_n\}$, $D = \{D_n\}$ and $E = \{E_n\}$ chain complexes and the short exact sequence:

$$0 \rightarrow C \xrightarrow{f} D \xrightarrow{g} E \rightarrow 0$$

for each p we get a induced exact sequence:

$$H_p(C) \xrightarrow{f_*} H_p(D) \xrightarrow{g_*} H_p(E).$$

Suppose now, that each line of the following diagram is exact and each square is commutative:

$$\begin{array}{ccccccc}
 & & \vdots & & \vdots & & \vdots \\
 & & \downarrow & & \downarrow & & \downarrow \\
 \cdots & \longrightarrow & C_{n+1} & \longrightarrow & D_{n+1} & \longrightarrow & E_{n+1} \longrightarrow \cdots \\
 & & \downarrow & & \downarrow & & \downarrow \\
 \cdots & \longrightarrow & C_n & \longrightarrow & D_n & \longrightarrow & E_n \longrightarrow \cdots \\
 & & \downarrow & & \downarrow & & \downarrow \\
 \cdots & \longrightarrow & C_{n-1} & \longrightarrow & D_{n-1} & \longrightarrow & E_{n-1} \longrightarrow \cdots \\
 & & \downarrow & & \downarrow & & \downarrow \\
 & & \vdots & & \vdots & & \vdots
 \end{array}$$

Then we can show that there exists a well defined homomorphism $\Delta: H_n(E) \rightarrow H_{n-1}(C)$ called the **connecting homomorphism**.

Theorem A.1.2. *If $0 \rightarrow C \xrightarrow{f} D \xrightarrow{g} E \rightarrow 0$ is a short exact sequence, then sequence:*

$$\cdots \xrightarrow{f_*} H_n(D) \xrightarrow{g_*} H_n(E) \xrightarrow{\Delta} H_{n-1}(C) \xrightarrow{f_*} \cdots$$

is exact.

Mayer Vietoris Sequence

For each cover $\mathcal{U}$ of X such that $\text{int}\mathcal{U}$ is also a cover for X , we denote $S_n^{\mathcal{U}}(X)$ the subgroup of $S_n(X)$ formed by those elements $\varphi: \sigma_n \rightarrow X$ with $\varphi(\sigma_p) \subset U$ for some $U \in \mathcal{U}$. Since for each i we have that $\text{im}(\partial_i \varphi) \subset \text{im}(\varphi)$ we can restrict ∂ to $\partial: S_n^{\mathcal{U}}(X) \rightarrow S_{n-1}^{\mathcal{U}}(X)$. Thus we can associate to each cover $\mathcal{U}$ of X a chain complex $S_*^{\mathcal{U}}(X)$ with the natural inclusion $i: S_*^{\mathcal{U}}(X) \rightarrow S_*(X)$, which is a chain map.

Note also that if $\mathcal{V}$ is a cover for Y and $f: X \rightarrow Y$ a map such that for each U in $\mathcal{U}$ $f(U) \subset V$ for some V in $\mathcal{V}$, then the induced map $f_{\#}: S_*^{\mathcal{U}} \rightarrow S_*^{\mathcal{V}}$ satisfies $f_{\#} \circ i_X = i_Y \circ f_{\#}$.

Theorem A.1.3. *If $\mathcal{U}$ covers X and $\text{int}\mathcal{U}$ also covers X , then $i_*: H_n(S_*^{\mathcal{U}}(X)) \rightarrow H_n(X)$ is an isomorphism for each n .*

If we consider such a cover $\mathcal{U}$ formed by only two open sets U and V we can construct a long exact sequence,

$$\cdots \xrightarrow{\Delta} H_n(U \cap V) \xrightarrow{g_*} H_n(S_*(U) \oplus S_*(V)) \xrightarrow{h_*} H_n(S_n^{\mathcal{U}}(X)) \xrightarrow{\Delta} H_{n-1}(U \cap V) \rightarrow \cdots$$

where g_* and h_* are induced by the following maps:

$$\begin{aligned} h: S_n(U) \oplus S_n(V) &\rightarrow S_n^{\mathcal{U}}(U \cup V) & g: S_n(U \cap V) &\rightarrow S_n(U) \oplus S_n(V) \\ (a'_i, a''_j) &\mapsto a'_i + a''_j & b_i &\mapsto (b_i, -b_i) \end{aligned}$$

and in view of Theorem A.1.3 we get the following long exact sequence called the **Mayer-Vietoris Sequence**:

$$\cdots \xrightarrow{\Delta} H_n(U \cap V) \xrightarrow{g_*} H_n(U) \oplus H_n(V) \xrightarrow{h_*} H_n(X) \xrightarrow{\Delta} H_{n-1}(U \cap V) \rightarrow \cdots$$

We are now able to compute some important results that will play important role later on. We shall do these computations by means of examples.

Example. Let $X = S^1$, we denote by z and z' the north and south poles of S^1 and x and y the points on the equator. Taking $U = S^1 - \{z'\}$ and $V = S^1 - \{z\}$ we can construct a Mayer-Vietoris sequence with the cover $\mathcal{U} = \{U, V\}$, but first note that U and V are contractible, i.e., have the same homotopy type as a point, hence $H_1(U) = H_1(V) = 0$. Therefore, the Mayer-Vietoris sequence yields:

$$0 \rightarrow H_1(S^1) \xrightarrow{\Delta} \mathbb{Z} \oplus \mathbb{Z} \xrightarrow{g_*} \mathbb{Z} \oplus \mathbb{Z}$$

then Δ is a monomorphism hence $H_1(S^1)$ is isomorphic to $\text{im}(\Delta) = \ker(g_*)$. Now, an element in $H_0(U \cap V) \simeq \mathbb{Z} \oplus \mathbb{Z}$ is of the form $ax + by$ with $a, b \in \mathbb{Z}$, and by definition of g_* we get:

$$g_*(ax + by) = (i_*(ax + by), -j_*(ax + by))$$

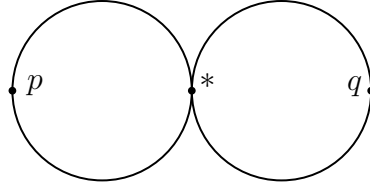
since U and V are path connected, $i_*(ax + by) = 0 \Leftrightarrow a = -b$ and similarly to $j_*(ax + by)$. Thus $\ker(g_*)$ is the subgroup of $H_0(U \cap V)$ consisting of all elements of the form $ax - ay = a(x - y)$. This is an infinite cyclic group generated by $x - y$, thus, isomorphic to $\mathbb{Z}$, hence:

$$H_1(S^1) = \mathbb{Z}.$$

For $n > 1$ we get that $H_n(U) = H_n(V) = H(U \cap V) = 0$, hence, the Mayer-Vietoris sequence gives us $H_n(S^1) = 0$.

Example. Consider $X = S^n$, z and z' the north and south poles respectively, and S^{n-1} the equator. By stereographic projection we know that $S^n - \{z\} \cong S^n - \{z'\} \cong \mathbb{R}^n$. Furthermore, $S^n - \{z, z'\} \cong \mathbb{R}^n - \{0\}$. Now take the cover $U = S^n - \{z\}$, $V = S^n - \{z'\}$, then $U \cap V = S^n - \{z, z'\}$, hence we get the following sequence from the Mayer-Vietoris sequence:

$$H_m(\mathbb{R}^n) \oplus H_m(\mathbb{R}^n) \xrightarrow{h_*} H_m(S^1) \xrightarrow{\Delta} H_m(S^{n-1}) \xrightarrow{g_*} H_{m-1}(\mathbb{R}^n) \oplus H_{m-1}(\mathbb{R}^n)$$



then for $m > 1$, we get $H_m(\mathbb{R}^n) = H_{m-1}(\mathbb{R}^n) = 0$ which implies that Δ is an isomorphism, i.e. $H_m(S^m) \simeq H_{m-1}(S^{n-1})$.

For $m = 1$ and $n > 1$ we get:

$$0 \rightarrow H_1(S^n) \xrightarrow{\Delta} \mathbb{Z} \xrightarrow{g_*} \mathbb{Z} \oplus \mathbb{Z}$$

but since $g_*(x) = (x, -x)$, we have that $g_*(x) = 0 \Leftrightarrow x = 0$, thus $\ker(g_*) = \{0\} = \text{im}(\Delta)$ and since in the previous sequence Δ is a monomorphism, it follows that $H_1(S^n) = 0$.

On the n -th level we get $H_n(S^n) \simeq H(S^{n-1}) \simeq \mathbb{Z}$. Thus by the previous example and using induction we have proved that:

$$H_m(S^n) \simeq \begin{cases} \mathbb{Z}, & \text{if } m = 0, n \\ 0, & \text{otherwise} \end{cases}.$$

Example Consider X the wedge product of two copies of S^1 , i.e., $S^1 \vee S^1$. If $n > 1$, from Mayer-Vietoris for the cover $U = S^1 \vee S^1 - \{p\}$ and $V = S^1 \vee S^1 - \{q\}$, which implies $A \cap B = \{*\}$ as in the picture

we get

$$0 \rightarrow 0 \rightarrow H_n(X) \xrightarrow{\Delta} 0$$

hence $H_n(X) = 0$ for all $n > 1$.

For $n = 1$ we get

$$0 \rightarrow \mathbb{Z} \oplus \mathbb{Z} \xrightarrow{f} H_1(X) \xrightarrow{\Delta} \mathbb{Z} \xrightarrow{g} \mathbb{Z} \oplus \mathbb{Z}$$

and since $g(1) = (1, 1)$ hence g is injective, therefore its kernel is zero, which implies $\text{im}(\Delta) = 0 \Rightarrow \Delta \equiv 0$. Thus, $\ker(\Delta) = H_1(S^1 \vee S^1) = \text{im}(f)$. Now we just need to check what is $\text{im}(f)$. Since $\ker(f) = \{0\}$, f is a monomorphism hence $\mathbb{Z} \oplus \mathbb{Z}$ is isomorphic to $\text{im}(f)$, which concludes the result:

$$H_n(S^1 \vee S^1) \simeq \begin{cases} \mathbb{Z}, & \text{if } n=0 \\ \mathbb{Z} \oplus \mathbb{Z}, & \text{if } n = 1 \\ 0, & \text{otherwise} \end{cases}.$$

Example. Now, consider the torus $X = \mathbb{T}^2$. Let U be $\mathbb{T}^2$ with a small disk D removed from it, and V a disk slightly bigger than D containing D . Note that U is homotopy equivalent to $S^1 \vee S^1$, V homotopy equivalent to a point and $U \cap V$ homotopy equivalent to S^1 . Take $n \geq 3$, the Mayer-Vietoris sequence gives us:

$$0 \rightarrow 0 \rightarrow H_n(\mathbb{T}^2) \rightarrow 0$$

hence, $H_n(\mathbb{T}^2) = 0$ for $n \geq 3$.

If $n = 2$ the Mayer-Vietoris sequence yields:

$$0 \xrightarrow{h_*} H_2(\mathbb{T}^2) \xrightarrow{\Delta} \mathbb{Z} \xrightarrow{g_*} (\mathbb{Z} \oplus \mathbb{Z}) \oplus 0$$

thus Δ is a monomorphism and by the definition of $g_* = (i_*, -j_*)$ with i and j inclusions, we get that $j_* \equiv 0$ (note that the j_* is defined from $\mathbb{Z}$ to $H_2(V) = 0$). Then since i_* is induced by inclusion, taking 1 a generator for $H_1(U \cap V) \simeq \mathbb{Z}$ we get $i_*(1) = a + b - a - b \in H_1(U) \simeq \mathbb{Z} \oplus \mathbb{Z}$, thus $i_*(1) = 0 \Rightarrow i_* \equiv 0$. Hence, $\text{im}(\Delta) = \ker(g_*) = \mathbb{Z}$, hence $H_2(\mathbb{T}) \simeq \mathbb{Z}$.

For $n = 1$ the Mayer-Vietoris sequence gives us:

$$\mathbb{Z} \xrightarrow{g_*^1} \mathbb{Z} \oplus \mathbb{Z} \oplus 0 \xrightarrow{h_*} H_1(\mathbb{T}^2) \xrightarrow{\Delta} \mathbb{Z} \xrightarrow{g_*^2} \mathbb{Z} \oplus \mathbb{Z}$$

then, from what we've done so far, $g_*^1 \equiv 0 \Rightarrow \text{im}(g_*^1) = 0 = \ker(h_*)$ therefore h_* is a monomorphism into $H_1(\mathbb{T}^2)$. Now, recall that $g_*^2(1) = (1, 1, 0)$, so g_*^2 is injective, then $\text{im}(\Delta) = \ker(g_*^2) = \{0\}$, which implies that $\ker(\Delta) \simeq H_1(\mathbb{T}^2)$, but $\text{im}(g_1^*) = \ker(\Delta)$ and $\text{im}(g_1^*) = \mathbb{Z} \oplus \mathbb{Z}$, hence, $H_1(\mathbb{T}^2) = \mathbb{Z} \oplus \mathbb{Z}$, and we have the homology groups of the torus:

$$H_n(\mathbb{T}^2) \simeq \begin{cases} \mathbb{Z}, & \text{if } n = 0, 2 \\ \mathbb{Z} \oplus \mathbb{Z}, & \text{if } n = 1 \\ 0, & \text{otherwise} \end{cases}.$$

Example. Consider now the genus 2 surface Σ_2 . The cover can be take as follows: consider a copy of S^1 in Σ_2 going between the two hole of the surface, in such a way that if we were to cut the surface along this S^1 we would get two copies of a torus with a disk removed, each disk having boundary exactly this copy of S^1 . Now consider U the part on the surface containing one the holes, the constructed S^1 , and a small part of the other part of the surface. Define V as a similar set, but containing the other part of the surface. Of course, by this construction we get $U \simeq S^1 \vee S^1 \simeq_H V$ and $U \cap V \simeq_H S^1$ (here $\simeq_H$ just means that they have the same homotopy type). For $n \geq 3$ we get the following from Mayer-Vietoris:

$$0 \rightarrow 0 \rightarrow H_n(\Sigma_2) \xrightarrow{\Delta} 0$$

hence $H_n(\Sigma_2) = 0$ for $n \geq 3$.

For $n = 2$ we get:

$$0 \rightarrow 0 \rightarrow H_2(\Sigma_2) \xrightarrow{\Delta} \xrightarrow{g_*} (\mathbb{Z} \oplus \mathbb{Z}) \oplus (\mathbb{Z} \oplus \mathbb{Z})$$

then Δ is a monomorphism, hence $H_2(\Sigma_2) \simeq \text{im}(\Delta) = \ker(g_*)$. As we done in previous examples, we get that since g_* is induced by inclusion $g_* \equiv 0$, hence $\ker(g_*) = \mathbb{Z}$, thus $H_2(\Sigma_2) \simeq \mathbb{Z}$.

For $n = 1$:

$$\mathbb{Z} \xrightarrow{f} \mathbb{Z} \oplus \mathbb{Z} \oplus \mathbb{Z} \oplus \mathbb{Z} \xrightarrow{g} H_1(\Sigma_2) \xrightarrow{\Delta} \mathbb{Z} \xrightarrow{h} \mathbb{Z} \oplus \mathbb{Z}$$

Again, using similar arguments as before, we get $f \equiv 0$ and $h(1) = (1, -1) \Rightarrow h$ is injective $\Rightarrow \ker(h) = \{0\} = \text{im}(\Delta) \Rightarrow \ker(\Delta) = H_1(\Sigma_2)$, since $\ker(\Delta) = \text{im}(g) \Rightarrow \text{im}(g) = H_1(\Sigma_2)$ therefore $H_1(\Sigma_2) = \mathbb{Z} \oplus \mathbb{Z} \oplus \mathbb{Z} \oplus \mathbb{Z}$.

The Tor Functor

Let A be a group, if we take a short exact sequence

$$0 \rightarrow G_1 \xrightarrow{j} G_0 \xrightarrow{\tau} A \rightarrow 0$$

such that both G_1 and G_2 are free abelian groups, then we call this particular kind of exact sequence a **free resolution** for A . We know that tensoring this sequence by some group D , we do not guarantee that the sequence is exact on left anymore, and we get the following new right exact sequence:

$$G_1 \otimes D \xrightarrow{j \otimes \text{Id}} G_0 \otimes D \xrightarrow{\tau \otimes \text{Id}} A \otimes D \rightarrow 0$$

we then define $\text{Tor}(A, D) = \ker(j \otimes \text{Id})$. The Tor functor, in a way, measures how far from being a short exact sequence the above sequence is, in other words, it measures the extent to which $j \otimes \text{Id}$ fails to be a monomorphism.

Example. Consider the following sequence:

$$0 \rightarrow \mathbb{Z} \xrightarrow{j} \mathbb{Z} \xrightarrow{\tau} \mathbb{Z}_p \rightarrow 0$$

where j is multiplication by p , then $\text{im}(j) = p\mathbb{Z} = \ker(\tau)$, and τ is just the canonical projection. Since $\mathbb{Z}$ is a free abelian group, then the sequence is a free resolution for $\mathbb{Z}_p$. Now, tensoring it by $\mathbb{Z}_q$ we get

$$\mathbb{Z} \otimes \mathbb{Z}_q \xrightarrow{j \otimes \text{Id}} \mathbb{Z} \otimes \mathbb{Z}_q \xrightarrow{\tau \otimes \text{Id}} \mathbb{Z}_p \otimes \mathbb{Z}_q \rightarrow 0$$

and since $\mathbb{Z} \otimes \mathbb{Z}_q = \mathbb{Z}_q$, we get $j \otimes \text{Id} = j$, therefore $\ker(j)$ is just the set of elements of order q in $p\mathbb{Z}$, hence, $\text{Tor}(\mathbb{Z}_p, \mathbb{Z}_q) = \ker(j) = \mathbb{Z}_{(p,q)}$ where $(p, q) = \gcd(p, q)$.

The Tor functor satisfies the following important properties:

- If A is free abelian, then $\text{Tor}(A, B) = 0$ for any abelian group B ;
- $\text{Tor}(A, B)$ is independent of the resolution of A ;
- $\text{Tor}(A, B) \simeq \text{Tor}(B, A)$;
- $\text{Tor}(A \oplus B, C) \simeq \text{Tor}(A, C) \oplus \text{Tor}(B, C)$;
- $\text{Tor}(A, B \oplus C) \simeq \text{Tor}(A, B) \oplus \text{Tor}(A, C)$.

Pairs of Spaces

We understand a **pair of spaces** as a pair (X, A) where X is a topological space and A is a subspace of X with the induced topology. If (X, A) is a pair of spaces, we can define $S_*(A)$ as a subcomplex of $S_*(X)$. The singular chain complex of X mod A is then defined as

$$S_*(X, A) = \frac{S_*(X)}{S_*(A)}$$

we can, of course, define the homology of this chain complex, the so called **relative singular homology** of X mod A , defined as

$$H_n(X, A) = H_n \left(\frac{S_*(X)}{S_*(A)} \right).$$

Any pair of spaces has an exact homology sequence given by

$$\cdots \rightarrow H_n(A) \xrightarrow{i_*} H_n(X) \xrightarrow{\pi_*} H_n(X, A) \xrightarrow{\Delta} H_{n-1}(A) \rightarrow \cdots$$

where i_* is induced by inclusion and π_* induced by the canonical projection. In some sense, $H_*(X, A)$ measures how far $i_*: H_n(A) \rightarrow H_n(X)$ is from being an isomorphism, in fact i_* is an isomorphism if and only if $H_n(X, A) = 0$.

Now, fixing an abelian group G for each pair of spaces (X, A) we can use $S_*(X, A)$ to construct a new complex $S_*(X, A; G) = S_*(X, A) \otimes G$, called **singular chain complex of (X, A) with coefficients in G** . The homology of $S_*(X, A; G)$ is denoted by $H_*(X, A; G)$.

The following theorem will be strongly utilized in the future, when dealing with homologies of manifolds, it is called the **universal coefficients theorem** for homology, there exists also a version of this theorem for cohomology, as we shall see:

Theorem A.1.4 (U.C.T.H). *For any pair of spaces (X, A) and G an abelian group:*

$$H_n(X, A; G) \simeq H_n(X, A) \otimes G \oplus \text{Tor}(H_{n-1}(X, A), G).$$

this theorem allows us to compute homology of spaces (and pairs of spaces) in other coefficient groups not only $\mathbb{Z}$.

Cohomology

Take $\text{Hom}(A, G)$ as usual for A and G abelian groups. If B is an abelian group and $\varphi: A \rightarrow B$ is a homomorphism, then there is an induced homomorphism:

$$\varphi^\#: \text{Hom}(B, G) \rightarrow \text{Hom}(A, G)$$

defined by $\varphi^\#(f) = f \circ \varphi$. Note that if $\psi: B \rightarrow C$, then

$$(\psi \circ \varphi)^\# = \varphi^\# \circ \psi^\#.$$

For a chain complex $\{C_n, \partial\}$ and an abelian group G , define:

$$C^n = \text{Hom}(C_n, G)$$

then the boundary operator $\partial: C_n \rightarrow C_{n-1}$ induces:

$$\partial^\#: C^n \rightarrow C^{n+1}$$

satisfying, obviously $\partial^\# \circ \partial^\# = 0$. This leads us to define a **cochain complex** $\{C^n, \delta\}$ where $\delta: C^n \rightarrow C^{n+1}$ and $\delta \circ \delta = 0$. The δ operator is called **coboundary operator**.

If $\{C^n, \delta\}$ and $\{D^n, \delta'\}$ are two cochain complexes, then a cochain map f of degree k is a collection

$$f: C^n \rightarrow D^{n+k}$$

such that $f \circ \delta = \delta' \circ f$.

Two cochain maps are said to be **cochain homotopic** if there exists a collection of homomorphisms $T: C^n \rightarrow D^{n-1}$ such that:

$$\delta' \circ T + T \circ \delta = f - g$$

T is called a **cochain homotopy**.

Note that if $\{C_n, \partial\}$ and $\{D_n, \partial'\}$ are chain complexes and $f, g: \{C_n, \partial\} \rightarrow \{D_n, \partial'\}$ are cochain homotopic, then for any abelian group G :

$$f^\#, g^\#: \{\text{Hom}(D_n, G), \partial'^\#\} \rightarrow \{\text{Hom}(C_n, G), \partial^\#\}$$

are cochain homotopic.

Let $C = \{C^n, \delta\}$ be a cochain complex. Define $Z^n(C) = \ker(\delta: C^n \rightarrow C^{n+1})$ the group of n -**cocycles** and $B^n(C) = \text{im}(\delta: C^{n-1} \rightarrow C^n)$ the group of n -**coboundaries**, then we can define the n -th **cohomology group** of C as:

$$H^n(C) = \frac{Z^n(C)}{B^n(C)}.$$

If $0 \rightarrow A \rightarrow B \rightarrow C \rightarrow 0$ is a short exact sequence of cochain maps, then there exists a long exact sequence:

$$\cdots \rightarrow H_n(A) \rightarrow H_n(B) \rightarrow H_n(C) \xrightarrow{\Delta} H^{n+1}(A) \rightarrow \cdots$$

Let (X, A) be a pair of spaces and G an abelian group. Let

$$\delta: S^n(X, A; G) \rightarrow S^{n+1}(X, A; G)$$

be given by $\delta - \partial^\#$ where $S^n(X, A; G) = \text{Hom}(S_n(X, A), G)$, this defines the singular cochain complex of the pair (X, A) where the cohomology is given by:

$$H^*(X, A; G).$$

It is important to note here that every covariant property of homology becomes a contravariant property in cohomology, hence the study of cohomology is a dual study to that of homology.

In particular, if $f: (X, A) \rightarrow (Y, B)$ is a map of pairs, then there is an induced homomorphism

$$f^*: H^*(Y, B; G) \rightarrow H^*(X, A; G).$$

If $g: (Y, B) \rightarrow (X, A)$ is another map of pairs, then the induced homomorphisms satisfy $(g \circ f)^* = f^* \circ g^*$.

The Ext Functor

Recall that a short exact sequence $0 \rightarrow A \rightarrow B \rightarrow C \rightarrow 0$ is called **split exact sequence** when B can be written as a direct sum of A and C . The split exact sequence has other characterizations but we will return to them if necessary.

Proposition A.1.1. *If $0 \rightarrow F \xrightarrow{i} H \xrightarrow{\pi} K \rightarrow 0$ is a split exact sequence, and G is abelian, then*

$$0 \rightarrow \operatorname{Hom}(K, G) \xrightarrow{\pi^\#} \operatorname{Hom}(H, G) \xrightarrow{i^\#} \operatorname{Hom}(F, G) \rightarrow 0$$

is exact (not necessarily split).

Note that in the Proposition A.1.1 we have that the initial sequence must be split so that the final sequence is exact. If the former is not split, then the latter need not be exact. As we did previously with the Tor functor, we can measure this failure of being exact by using a new functor, the Ext functor.

Let E be an abelian group and take a free resolution for E :

$$0 \rightarrow R \xrightarrow{i} F \xrightarrow{\pi} E \rightarrow 0$$

then for an abelian group G , take the sequence:

$$0 \rightarrow \operatorname{Hom}(E, G) \xrightarrow{\pi^\#} \operatorname{Hom}(F, G) \xrightarrow{i^\#} \operatorname{Hom}(R, G)$$

this sequence is left exact, but not exact on the right. Thus to measure how far from being exact, we use the functor:

$$\operatorname{Ext}(E, G) = \operatorname{coker}(i^\#) = \frac{\operatorname{Hom}(R, G)}{\operatorname{im}(i^\#)}.$$

The Ext functor satisfies the following important properties:

- If E is free abelian, then $\operatorname{Ext}(E, G) = 0$;
- $\operatorname{Ext}(E, G)$ is independent of choice of resolution for E ;
- Given $f: E \rightarrow E'$ and $g: G \rightarrow G'$ there are induced homomorphisms $f^*: \operatorname{Ext}(E', G) \rightarrow \operatorname{Ext}(E, G)$ and $g^*: \operatorname{Ext}(E, G) \rightarrow \operatorname{Ext}(E, G')$. In other words, $\operatorname{Ext}(E, G)$ is covariant in G and contravariant in E ;
- If $0 \rightarrow A \xrightarrow{j} B \xrightarrow{k} C \rightarrow 0$ is short exact, then

$$0 \rightarrow \operatorname{Hom}(C, G) \xrightarrow{k^\#} \operatorname{Hom}(B, G) \xrightarrow{j^\#} \operatorname{Hom}(A, G) \rightarrow \operatorname{Ext}(C, G) \xrightarrow{k^*} \operatorname{Ext}(B, G) \xrightarrow{j^*} \operatorname{Ext}(A, G) \rightarrow 0$$

is exact.

Products

We shall now establish a useful notation, note that an element of a cochain is a function from the respective chain complex to an abelian group G , hence, we can instead of using the usual notation $\varphi(c)$ for $\varphi \in S^n(X, A; G)$ and $c \in S_n(X, A)$, use the notation of a pairing, so:

$$\varphi(c) = \langle \varphi, c \rangle.$$

Furthermore, if $f: (X, A) \rightarrow (Y, B)$ is a map of pairs, we have the following property of such a pairing:

$$\langle \varphi, f_{\#}(c) \rangle = \langle f^{\#}(\varphi), c \rangle$$

which of course implies $\langle \delta\varphi, c \rangle = \langle \varphi, \partial c \rangle$. In this notation, we may say that a chain φ is a cocycle if and only if $\langle \delta\varphi, c \rangle = 0$ for every $c \in S_{n+1}(X, A)$, or equivalently, $\langle \varphi, \partial c \rangle = 0$ for every $c \in S_{n+1}(X, A)$. Thus, φ is a cocycle if, and only if $\varphi(Z_n(X, A)) = 0$.

Now, take $x \in H^n(X, A; G)$ represented by the cocycle φ and $y \in H_n(X, A)$ represented by the cycle c , we define

$$\begin{aligned} \langle \bullet, \bullet \rangle: H^n(X, A; G) \otimes H_n(X, A) &\rightarrow G \\ x \otimes y &\mapsto \langle x, y \rangle = \langle \varphi, c \rangle \end{aligned}$$

which is called the **Kronecker Index**, and may also be viewed as a homomorphism:

$$\alpha: H^n(X, A; G) \rightarrow \text{Hom}(H_n(X, A), G).$$

Proposition A.1.2. *If X is a topological space and $\{X_{\alpha}\}_{\alpha \in \Lambda}$ is its decomposition into path components, then:*

$$H^0(X; G) \simeq \prod_{\alpha \in \Lambda} G_{\alpha}$$

'the direct product of copies of G , one for each component of X .

Theorem A.1.5 (Excision). *Let (X, A) be a pair of spaces and $U \subset A$ a subset with $\bar{U} \subseteq \text{int}(A)$. Then the inclusion map of pairs:*

$$i: (X - U, A - U) \rightarrow (X, A)$$

induces an isomorphism

$$i^*: H^*(X, A; G) \rightarrow H^*(X - U, A - U; G).$$

The following, is the already announced **Universal Coefficients Theorem for Cohomology**, it will also be strongly used on previous studies of homology and cohomology of manifolds.

Theorem A.1.6 (U.C.T.C.). *Given a particular pair (X, A) of spaces and an abelian group G , there exists a split exact sequence:*

$$0 \rightarrow \text{Ext}(H_{n-1}(X, A), G) \rightarrow H^n(X, A; G) \xrightarrow{\alpha} \text{Hom}(H_n(X, A), G) \rightarrow 0$$

And from that, we obviously get:

$$H^n(X, A; G) \simeq \text{Hom}(H_n(X, A), G) \oplus \text{Ext}(H_{n-1}(X, A), G).$$

It is interesting to note that, contrary to our intuition the cohomology isn't just the group of homomorphisms defined on the homology into an abelian group G , but it also has an extra term on it $\text{Ext}(H_{n-1}(X, A), G)$. As we have seen before, if $H_{n-1}(X, A)$ is a

free abelian group, then we are allowed to use our intuition to describe the cohomology group in terms of homology.

Take $\{C_n, \partial\}$ and $\{D_n, \partial\}$ two chain complexes. Define:

$$(C \otimes D)_n = \sum_k C_k \otimes D_{n-k}$$

with the boundary operator

$$\partial: C_p \otimes D_q \rightarrow C_{p-1} \otimes D_q \oplus C_p \otimes D_{q-1}$$

given by

$$\partial(c \otimes d) = \partial c \otimes d + (-1)^p c \otimes \partial d,$$

then $\{(C \otimes D)_n, \partial\}$ is also a chain complex, it is easy to verify that $\partial \circ \partial = 0$.

Now, let C be a free chain complex, and the short exact sequence:

$$0 \rightarrow Z_n(C) \xrightarrow{i} C_n \xrightarrow{\partial} B_n(C) \rightarrow 0$$

with i being the inclusion, this exact sequence actually splits, and as mentioned before, one of the characterizations that split exact sequences have is that there exists a map $\varphi: C_n \rightarrow Z_n(C)$ such that $\varphi \circ i = \text{Id}_{Z_n(C)}$. Denoting by $\varphi = \pi \circ \varphi$ the composition

$$C_n \xrightarrow{\varphi} Z_n(C) \xrightarrow{\pi} H^n(C)$$

where π is the canonical projection, then we have the following theorem:

Theorem A.1.7. *If C and D are free chain complexes then the chain map*

$$\Phi \otimes \text{Id}: C \otimes D \rightarrow H_*(C) \otimes D$$

induces an isomorphism

$$(\Phi \otimes \text{Id})_*: H_n(C \otimes D) \rightarrow H_n(H_*(C) \otimes D).$$

The following is another important theorem that will help us on our objective of studying manifolds, it is called the **Künneth's Formula**:

Theorem A.1.8. *If C and D are free chain complexes, then*

$$H_n(C \otimes D) \simeq \bigoplus_{p+q=n} H_p(C) \otimes H_q(D) \oplus \bigoplus_{r+s=n-1} \text{Tor}(H_r(C), H_s(D)).$$

Given spaces X and Y , the previous theorem allows us to write:

$$H_n(S_*(X) \otimes S_*(Y)) \simeq \bigoplus_{p+q=n} H_p(X) \otimes H_q(Y) \oplus \bigoplus_{r+s=n-1} \text{Tor}(H_r(X), H_s(Y))$$

our objective now is to relate $H_n(S_*(X) \otimes S_*(Y))$ to $H_n(X \times Y)$. Without showing the details, because that would take more time than necessary, we can say that using a little of category theory, precisely using a theorem called *acyclic models theorem*, and the Künneth formula, we can prove the following result that establishes such a relation between the groups that we wanted:

Theorem A.1.9 (Eilenberg-Zilber). *For any spaces X and Y and any integer k there is an isomorphism*

$$\phi_*: H_k(X \times Y) \rightarrow H_k(S_*(X) \otimes S_*(Y)).$$

this allows us to rewrite the Künneth's formula in its full glory as:

$$H_n(X \times Y) \simeq \bigoplus_{p+q=n} H_p(X) \otimes H_q(Y) \oplus \bigoplus_{r+s=n-1} \text{Tor}(H_r(X), H_s(Y)).$$

Now, fix a chain map

$$\varphi: S_*(X \times Y) \rightarrow S_*(X) \otimes S_*(Y)$$

then the composition

$$H_p(X) \otimes H_q(Y) \rightarrow H_{p+q}(S_*(X) \otimes S_*(Y)) \xrightarrow{\varphi^{-1}} H_{p+q}(X \times Y)$$

where the first homomorphism takes $\{x\} \otimes \{y\}$ to $\{x \otimes y\}$ is called **homology external product**.

The image of $\{x\} \otimes \{y\}$ under this composition is usually denoted as $\{x\} \times \{y\}$. Note that, from Künneth's formula

$$0 \rightarrow \bigoplus_{p+q=n} H_p(X) \otimes H_q(Y) \xrightarrow{\times} H_n(X \times Y) \rightarrow \bigoplus_{r+s=n-1} \text{Tor}(H_r(X), H_s(Y)) \rightarrow 0$$

thus, $\times$ is a monomorphism.

Now, we want to give an analog for this product, but for cohomology, that is, a product

$$H^p(X; G_1) \otimes H^q(Y; G_2) \rightarrow H^{p+q}(X \times Y; G_1 \otimes G_2).$$

Take $\alpha \in S^p(X; G_1)$ and $\beta \in S^q(Y; G_2)$. Denote by $\alpha \times \beta$ the homomorphism given by

$$S_{p+q}(X \times Y) \xrightarrow{\varphi} S_*(X) \otimes S_*(Y) \xrightarrow{\alpha \otimes \beta} G_1 \otimes G_2$$

thus $\alpha \otimes \beta: S_{p+q}(X \times Y) \rightarrow G_1 \otimes G_2$, hence $\alpha \times \beta \in S^{p+q}(X \times Y; G_1 \otimes G_2)$, and this defines an external product on cochains. The following proposition is sometimes called the *derivation* rule for such a product, it tells us how to compute the boundary of the external product of two cochains:

Proposition A.1.3. *If $\alpha \in S^p(X; G_1)$ and $\beta \in S^q(Y; G_2)$ are cochains and $\alpha \times \beta \in S^{p+q}(X \times Y; G_1 \otimes G_2)$ is their external product, then*

$$\delta(\alpha \times \beta) = (\delta\alpha) \times \beta + (-1)^p \alpha \times \delta\beta.$$

as a corollary, we have that this product on cochains induces an external product on cohomology:

Corollary A.1.1. *The previous theorem induces a well defined external product on cohomology groups:*

$$H^p(X; G_1) \otimes H^q(Y; G_2) \rightarrow H^{p+q}(X \times Y; G_1 \otimes G_2)$$

given by $\{\alpha\} \otimes \{\beta\} = \{\alpha \times \beta\}$.

Note that taking R an associative commutative ring with unit, we can define $\mu: R \otimes R \rightarrow R$ given by $\mu(a \otimes b) = ab$. On the previous construction, taking $G_1 = G_2 = R$, we get the following external product on cohomology

$$\times_1: H^p(X; R) \otimes H^q(Y; R) \rightarrow H^{p+q}(X \times Y; R)$$

taking $\{\alpha\} \otimes \{\beta\}$ to $\{\alpha \times_1 \beta\}$.

Lemma A.1.2. $\{\alpha\} \in H^p(X; R)$, $\{\beta\} \in H^q(X; R)$. Define $T: X \times Y \rightarrow Y \times X$ by $T(x, y) = (y, x)$, then:

$$T^*: H^{p+q}(Y \times X; R) \rightarrow H^{p+q}(X \times Y; R)$$

has

$$T^*(\{\beta\} \times_1 \{\alpha\}) = (-1)^{pq}(\{\alpha\} \times_1 \{\beta\}).$$

Lemma A.1.3. $\{\alpha\} \in H^*(X; R)$, $\{\beta\} \in H^*(Y; R)$ and $f: X' \rightarrow X$, $g: Y' \rightarrow Y$, then

$$(f \times g)^*(\{\alpha\} \times_1 \{\beta\}) = f^*(\{\alpha\}) \times_1 g^*(\{\beta\}).$$

Lemma A.1.4. Taking $\{\alpha\}$ and $\{\beta\}$ as before, and $\{\gamma\} \in H^*(W; R)$, we have

$$(\{\alpha\} \times_1 \{\beta\}) \times_1 \{\gamma\} = \{\alpha\} \times_1 (\{\beta\} \times_1 \{\gamma\}).$$

Note that, $Y = \{p\}$ then $H^0(Y; R) \simeq R$, then

$$H^*(X; R) \otimes H^*(Y; R) \rightarrow H^*(X \times Y; R)$$

has the form

$$H^*(X; R) \otimes R \rightarrow H^*(X; R)$$

this gives the structure of a graded R -module to $H^*(X; R)$. Not only the structure of a R -module can be given to a cohomology group, but considering $d: X \rightarrow X \times X$ the diagonal mapping $d(x) = (x, x)$, then the composition

$$H^p(X; R) \otimes H^q(X; R) \rightarrow H^{p+q}(X \times X; R) \xrightarrow{d^*} H^{p+q}(X; R)$$

sending $\{\alpha\} \otimes \{\beta\}$ to $d^*(\{\alpha\} \times_1 \{\beta\})$ defines a multiplication in the R -module $H^*(X; R)$ called **cup product** written as $\{\alpha\} \smile \{\beta\}$. Hence

$$\{\alpha\} \smile \{\beta\} = d^*(\{\alpha\} \times_1 \{\beta\}).$$

Which gives to $H^*(X; R)$ a new structure, written as a theorem:

Theorem A.1.10. For R a commutative ring with unit, X a topological space, $H^*(X; R)$ is a commutative associative graded R -algebra with unit. Any continuous function $f: X' \rightarrow X$ induces an R -algebra homomorphism

$$F^*: H^*(X; R) \rightarrow H^*(X'; R)$$

of degree zero.

If we take $\varphi: \sigma_p \rightarrow X$ a singular p -simplex, we can define its **i -th front face** as

$${}_i\varphi(t_0, \dots, t_i) = \varphi(t_0, \dots, t_i, 0, \dots, 0)$$

as well as its **j -th back face** as

$$\varphi_j(t_0, \dots, t_i) = \varphi(0, \dots, 0, t_0, \dots, t_j).$$

Then, by using this notation we get the following property satisfied by the cup product:

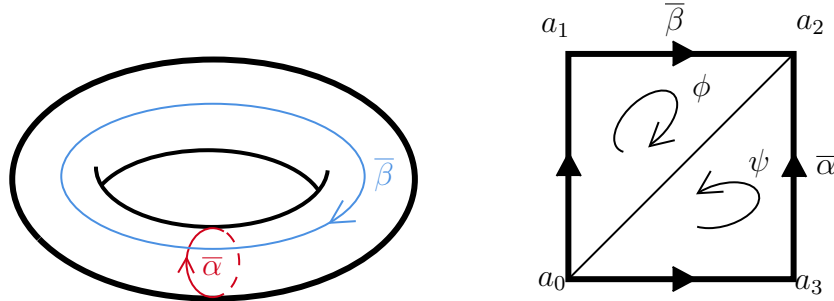
$$\langle \alpha \smile \beta, \varphi \rangle = \langle \alpha, {}_p\varphi \rangle \langle \beta, \varphi_p \rangle$$

more geometrically, we are defining the cup product of α and β on a singular p -simplex, as applying the cochain α to the i -th front face of φ and then β to the back j -th face of φ .

Example Let us compute the cohomology ring of $\mathbb{T}\mathbb{T}^2$. Recall that $H_1(\mathbb{T}^2; \mathbb{Z}) \simeq \mathbb{Z} \oplus \mathbb{Z}$ and take as generators of such a group the elements $\bar{\alpha}$ and $\bar{\beta}$. For $H_2(\mathbb{T}^2; \mathbb{Z}) \simeq \mathbb{Z}$ take as generator the 2-chain $\phi - \psi$ given by

$$\begin{array}{lll} \phi(0) = a_0 & \phi(1) = a_1 & \phi(2) = a_2 \\ \psi(0) = a_0 & \psi(1) = a_3 & \psi(2) = a_2 \end{array}$$

following the next figure:



By using the universal coefficients theorem for cohomology, we get $H^1(\mathbb{T}^2; \mathbb{Z}), \mathbb{Z}$, since $H_0(\mathbb{T}^2; \mathbb{Z}) \simeq \mathbb{Z}$ then the $\text{Ext}(H_0(\mathbb{T}^2; \mathbb{Z}), \mathbb{Z})$ part of such formula is zero. Hence, we may choose generators α and β such that

$$\begin{array}{ll} \alpha(\bar{\alpha}) = 1 & \alpha(\bar{\beta}) = 0 \\ \beta(\bar{\alpha}) = 0 & \beta(\bar{\beta}) = 1 \end{array}$$

Now, we get:

$$\begin{aligned} \langle \alpha \smile \beta, \phi - \psi \rangle &= \langle \alpha \smile \beta \rangle - \langle \alpha \smile \beta, \psi \rangle \\ &= \langle \alpha, {}_1\phi \rangle \langle \beta, \phi_1 \rangle - \langle \alpha, {}_1\psi \rangle \langle \beta, \psi_1 \rangle \\ &= \langle \alpha, \bar{\alpha} \rangle \langle \beta, \bar{\beta} \rangle - \langle \alpha, \bar{\beta} \rangle \langle \beta, \bar{\alpha} \rangle \\ &= 1 \cdot 1 - 0 \cdot 0 = 1 \end{aligned}$$

on the other hand:

$$\begin{aligned}
 \langle \alpha \smile \alpha, \phi - \psi \rangle &= \langle \alpha \smile \alpha \rangle - \langle \alpha \smile \alpha, \psi \rangle \\
 &= \langle \alpha, {}_1\phi \rangle \langle \alpha, \phi_1 \rangle - \langle \alpha, {}_1\psi \rangle \langle \alpha, \psi_1 \rangle \\
 &= \langle \alpha, \bar{\alpha} \rangle \langle \alpha, \bar{\beta} \rangle - \langle \alpha, \bar{\beta} \rangle \langle \alpha, \bar{\alpha} \rangle \\
 &= 1 \cdot 0 - 0 \cdot 1 = 0
 \end{aligned}$$

Similarly, we get $\langle \beta \smile \beta, \phi - \psi \rangle = 0$. Now, since $H^2(\mathbb{T}^2; \mathbb{Z}) \simeq \text{Hom}(H_2(\mathbb{T}^2; \mathbb{Z}), \mathbb{Z}) \oplus \text{Ext}(H_1(\mathbb{T}^2; \mathbb{Z}), \mathbb{Z})$ and $H_1(\mathbb{T}^2; \mathbb{Z})$ is free abelian, then $H^2(\mathbb{T}^2; \mathbb{Z}) \simeq \text{Hom}(H_2(\mathbb{T}^2; \mathbb{Z}), \mathbb{Z}) \simeq \mathbb{Z}$ since $H_2(\mathbb{T}^2; \mathbb{Z}) \simeq \mathbb{Z}$. Thus, $H^*(\mathbb{T}^2; \mathbb{Z})$ is the graded algebra over $\mathbb{Z}$ generated by the elements α and β of degree 1 subject to the relations

$$\alpha^2 = 0 \qquad \beta^2 = 0 \qquad \alpha\beta = -\beta\alpha.$$

Since $H^*(X; R)$ is a cohomology ring with the multiplication given by the cup product, this means that we can write this ring as

$$H^*(X; R) = \bigoplus_{k=1}^{\infty} H^k(X; R)$$

each level $H^k(X; R)$ is just the cohomology group as we know it.

Now, there is yet another kind of multiplication we can give to a cohomology group, but this time, we define a multiplication between an element of homology and an element of cohomology. Take X a topological space and R a commutative ring with unit. If we take $\alpha \in S^p(X; R)$ we know that α is a map from $S_p(X)$ to R . But setting $\alpha = 0$ in all elements of $S^*(X; R)$ with dimension different from p we can look at α as an element of $S^*(X; R)$.

The composition

$$S_*(X) \xrightarrow{\tau} S_*(X) \otimes S_*(X) \xrightarrow{\alpha \otimes \text{id}} R \otimes S_*(X)$$

where τ is a generalization of the diagonal map, when tensored with R yields:

$$R \otimes S_*(X) \xrightarrow{\otimes \tau} R \otimes S_*(X) \otimes S_*(X) \xrightarrow{\text{Id} \otimes \alpha \otimes \text{id}} R \otimes R \otimes S_*(X) \xrightarrow{\mu \otimes \text{Id}} R \otimes S_*(X).$$

Now, if $c \in S_n(X; R) = R \otimes S_n(X)$ we define the **cap product** of α and c as $\alpha \frown c$ by being the image of c under this composition. Note:

$$\alpha \frown x \in S_{n-p}(X; R).$$

Define now $\tau(\phi) = \sum_{i+j=n} i\phi \otimes \phi_j$, then τ is called **Alexander-Whitney diagonal approximation**. Consider this diagonal approximation and ϕ a singular n -simplex, then from the previous sequence we have

$$1 \otimes \phi \mapsto 1 \otimes \left\{ \sum_{i+j=n} i\phi \otimes \phi_j \right\} \mapsto 1 \otimes \alpha({}_p\phi) \otimes \phi_{n-p} \mapsto \alpha({}_p\phi) \otimes \phi_{n-p}$$

if we interpret $R \otimes S_*(X)$ as the free R -module generated by the singular simplexes of X , then

$$\alpha \frown \phi = \alpha({}_p\phi) \phi_{n-p}.$$

Now, note that if we take $\alpha \in S^p(X; R)$, $\beta \in S^q(X; R)$ and $\phi \in S_{p+q}(X; R)$ where ϕ is a singular simplex, then we have

$$\begin{aligned}\langle \beta \smile \alpha, \phi \rangle &= \langle \beta, {}_q\phi \rangle \langle \alpha, \phi_p \rangle \\ &= \beta({}_q\phi) \alpha(\phi_p)\end{aligned}$$

on the other hand

$$\langle \alpha, \beta \smile \phi \rangle = \langle \alpha, \beta({}_q\phi) \phi_p \rangle = \beta({}_q\phi) \alpha(\phi_p)$$

since this holds for any simplex ϕ , then for $c \in S_{p+q}(X; R)$ we get

$$\langle \beta \smile \alpha, c \rangle = \langle \alpha, \beta \smile c \rangle.$$

Finally, since $\alpha \smile c$ is a chain, we determine ∂ on $\alpha \smile c$. Take γ an arbitrary cochain, it is not hard to show that

$$\langle \gamma, \partial(\alpha \smile c) \rangle = \langle \gamma, (-1)^q(\alpha \smile \partial c - \delta \alpha \smile c) \rangle$$

it then follows that

$$(-1)^q \partial(\alpha \smile c) = (\alpha \smile \partial c) - (\delta \alpha \smile c).$$

Hence we can see that there is a well defined product on homology groups which takes the form:

$$H^q(X; R) \otimes H_n(X; R) \rightarrow H_{n-q}(X; R)$$

taking $\{\alpha\} \otimes \{c\}$ to $\{\alpha \smile c\}$.

To finalize the review on homology and cohomology, we establish the Künneth's formula on the context of cohomology:

Theorem A.1.11 (Cohomology Künneth's Formula). *Let G and G' be abelian groups with $\text{Tor}(G, G') = 0$. If $H_*(X; \mathbb{Z})$ and $H_*(Y; \mathbb{Z})$ are of finite type, then there is a split exact sequence*

$$0 \rightarrow \bigoplus_{p+q=n} H^p(X; G) \otimes H^q(Y; G') \xrightarrow{\simeq} H^n(X \times Y; G \otimes G') \rightarrow \bigoplus_{p+q=n+1} \text{Tor}(H^p(X; G), H^q(Y; G)) \rightarrow 0.$$

Hence, since the sequence splits, we can write:

$$H^n(X \times Y; G \otimes G') \simeq \bigoplus_{p+q=n} H^p(X; G) \otimes H^q(Y; G') \oplus \bigoplus_{p+q=n+1} \text{Tor}(H^p(X; G), H^q(Y; G)).$$

Hurevich's Theorem $n = 1$

The last result that we want to establish before starting or study on manifolds is one that will be strongly used to characterize the first homology groups of every connected manifold. This is the Hurevich's theorem for $n = 1$, the statement " $n = 1$ " comes from the face that there are versions of this theorem for other values of n as well, but for now we shall only use the case on the theorem.

Let us consider X a topological space based at a point x_0 . consider the map

$$h: \pi_1(X, x_0) \rightarrow H_1(X),$$

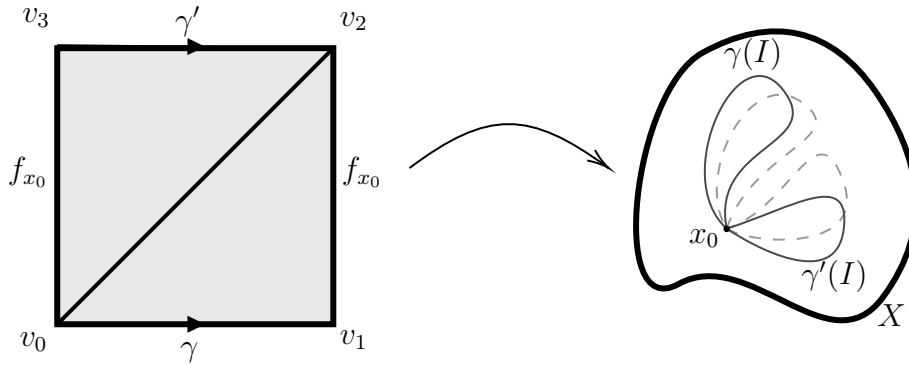
with $\pi_1(X, x_0)$ the fundamental group of X based on the point x_0 as usual. Take an element $[\gamma] \in \pi_1(X, x_0)$. By definition, we know that $\gamma: [0, 1] \rightarrow X$, thus we can think of γ as a singular 1-simplex in X , but since γ is a closed path, then γ is actually a 1-cycle in X , thus we simply consider $h([\gamma]) = \{\gamma\}$, where $\{\gamma\}$ denotes the homology class of the 1-cycle γ . Then we have the following theorem:

Theorem A.1.12 (Hurevich's $n = 1$). *If X is a topological space that is path connected and base pointed at x_0 , then h is a well defined homomorphism and the map $h': (\pi_1(X, x_0))_{Ab} \rightarrow H_1(X)$ is an isomorphism*

Proof. By considering the previous discussion about h we here have to prove the following facts:

- h is well defined;
- h is a homomorphism;
- h' is an isomorphism:
 - h is surjective;
 - $\ker(h) = \pi_1(X, x_0)_{Ab}$.

(h is well defined): If γ is homotopic to γ' via the homotopy $H: I \times I \rightarrow X$, for γ and γ' two representatives of the same class on $\pi_1(X, x_0)$, then let's divide $I \times I$ into the diagonal $\overline{v_0 v_3}$ following the picture below, and we get that H may be regarded as the sum of 2 singular 2-simplexes $\varphi_1: \sigma_2 \rightarrow X$ and $\varphi_2: \sigma_2 \rightarrow X$.



Now, note that

$$\partial\varphi_1 = f_{x_0} - D + \gamma$$

$$\partial\varphi_2 = \gamma' - D + f_{x_0}$$

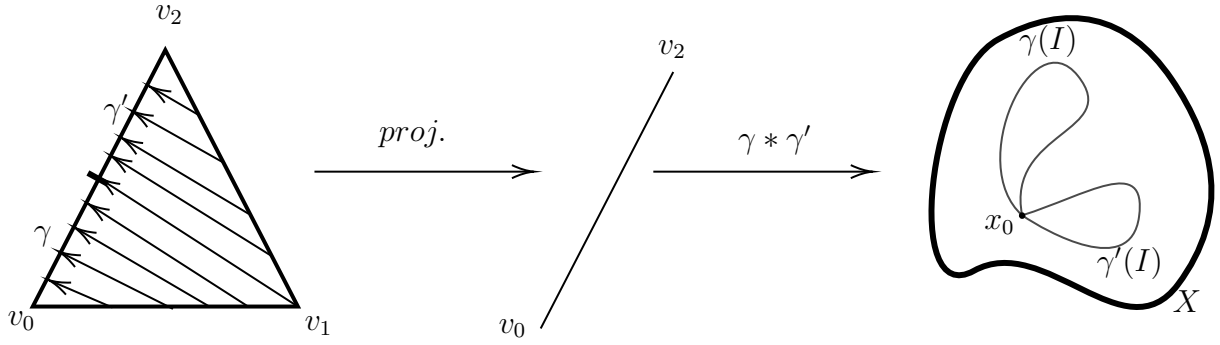
where D is the singular 1-simplex restricted to the diagonal on $I \times I$ and f_{x_0} is the constant map sending to x_0 . Now,

$$\partial(\varphi_1 - \varphi_2) = \partial\varphi_1 - \partial\varphi_2 = f_{x_0} - D + \gamma - \gamma' + D - f_{x_0} = \gamma - \gamma'$$

hence, $\gamma - \gamma' \in B_1(X)$, therefore $h([\gamma - \gamma']) = 0 \Rightarrow h([\gamma]) = h([\gamma'])$.

(h is a homomorphism): Let $[\gamma], [\gamma'] \in \pi_1(X, x_0)$ and write σ_2 as $[v_0, v_1, v_2]$ (forget the previous picture). Let us define a map $\psi: \sigma_2 \rightarrow X$ that is the composition of orthogonal projection of $[v_0, v_1, v_2]$ to $[v_0, v_2]$ and the product of γ and γ' (given by the group operation on $\pi_1(X, x_0)$)

$$[v_0, v_1, v_2] \xrightarrow{\text{proj.}} [v_0, v_2] \xrightarrow{\gamma * \gamma'} X$$



Now, $\partial\psi = \gamma' - \gamma * \gamma' + \gamma$, thus we get:

$$h([\gamma][\gamma']) = h([\gamma * \gamma']) = \{\gamma * \gamma'\} = \{\gamma + \gamma' + \partial\psi\}$$

since $\partial\psi \in B_1(X)$, we get that $h([\gamma][\gamma']) = \{\gamma\} + \{\gamma'\}$ thus h is a homomorphism.

(h is surjective): Let $\varphi = \sum n_i \varphi_i$ a 1-cycle in $S_1(X)$. We can split up the multiples of φ_i to make sure that every $n_i = \pm 1$. Now by possibly reversing the direction of φ_i we can make sure that every $n_i = 1$. Hence without loss of generality we can write $\varphi = \sum \varphi_i$.

If there is φ_i that is not a loop, there must exist some φ_j such that $\varphi_i + \varphi_j$ is a loop (since otherwise $\partial\varphi$ would not be trivial, but since we took φ as a cycle, $\partial\varphi = 0$). remember also that, $\varphi_i + \varphi_j$ is the same as $\varphi_i * \varphi_j$ since h is a homomorphism, then we can write $\varphi_i + \varphi_j$ as a single element $\varphi_i * \varphi_j$, finally, without loss of generality, we can take $\varphi = \sum \varphi_i$ with all φ_i being cycles.

Now, by hypothesis, X is path connected, therefore there is a path γ_i from x_0 to the basepoint of φ_i . Then by h , we have that $\gamma_i * \varphi_i \bar{\gamma}_i$ is homologous to $\gamma_i + \varphi_i + \bar{\gamma}_i$ and $\bar{\gamma}_i$ is homologous to $-\gamma_i$, which implies that $\gamma_i * \varphi_i \bar{\gamma}_i$ is homologous to φ_i since $H_1(X)$ is abelian.

Hence we may assume that the φ_i 's are all cycles based at x_0 , finally since $\sum \varphi_i$ is homologous to the product, we can take $\sum \varphi_i$ as a singular 1-simplex without changing the homology class ($\varphi_1 * \varphi_2 * \dots * \varphi_r$), in particular, this singular 1-simplex is a loop based at x_0 , which means that its homotopy class maps to the homology class of φ , as we wanted.

(Kernel): Let $\gamma * \gamma' * \bar{\gamma} * \bar{\gamma}'$ be the commutator of $\pi_1(X, x_0)$, then its image under h is $\gamma + \gamma' + \bar{\gamma} + \bar{\gamma}' \in H_1(X)$, which is of course zero in $H_1(X)$, hence $\pi_1(X, x_0)_{Ab} \subset \ker(h)$.

Now we need to prove the other inclusion. Suppose $[\gamma] \in \ker(h)$, we need to show that $[\gamma] = 0 \in \pi_1(X, x_0)_{Ab}$, since this means that $[\gamma]$ is in the commutator subgroup. Since $[\gamma] \in \pi_1(X, x_0)$ then γ is a loop, i.e., γ is a 1-cycle homologous to 0, which means that $\{\gamma\} \in B_1(X)$, thus there must exist some 2-cycle $\eta = \sum n_i \eta_i$ with $\partial\eta = \gamma$, and as before, we shall take every $n_i = \pm 1$.

Now, for each η_i let's write $\partial\eta_i = \tau_{i0} - \tau_{i1} + \tau_{i2}$ where τ_{ij} are 1-cycles. Note that:

$$\gamma = \partial \left(\sum n_i \eta_i \right) = \sum n_i \partial\eta_i = \sum n_i (\tau_{i0} - \tau_{i1} + \tau_{i2}) = \sum_{i,j} (-1)^j n_i \tau_{ij}.$$

But, we know that γ is a 1-cycle, hence τ_{ij} must cancel each other out in pairs, except for one, which equals γ . If we glue together the 2-simplexes, identifying the canceling pairs, we get a Δ -complex K .

Since the pairs that have been identified are actually the same map, the η_i together form a map $\eta: K \rightarrow X$. Let A be the 0-skeleton of K union the segment corresponding to

γ . We can slide the image of each vertex along a path from its original image to x_0 . This defines a homotopy of A with a new 0-skeleton which maps each point to x_0 and leaves γ unchanged. By the homotopy extension property since the 0-skeleton with the segment corresponding to γ is a subcomplex of X , we can extend this homotopy to a homotopy defined on all K .

If now we break the deformed K down in its simplexes, we get a new chain $\sum m_i \eta'_i$ but every singular 1-simplex τ'_{ij} on the boundary is now a loop at x_0 .

Basically what all this means is that we are deviding the 2 cycle into a sum of singular 2 simplexes, then fixing γ which is the boundary of this 2-simplex and then moving every vertex to x_0 , hence every singular 1-simplex is now a 1-cycle bounding a singular 2-simplex.

Now, since $\pi_1(X, x_0)_{Ab}$ is abelian, we can write $[\gamma] = \sum_{i,j} (-1)^j m_i [\tau'_{ij}]$ in the abelianization. We wrote $\partial \eta_i = \tau_{i0} - \tau_{i1} + \tau_{i2}$ but because h' is a homomorphism, the sums on j can be condensed to give $[\gamma] = \sum m_i [\partial \eta_i]$.

For each η_i we can deform the image of η_i to the image of one vertex by sliding the image of one edge through the image of its interior. This deformation is a homotopy between $\partial \eta_i$ and the constant map. Hence $[\partial \eta_i] = 0$, so $[f] = 0$ in the abelianization of $\pi_1(X, x_0)$, which completes the proof. $\square$

Manifolds

Let's begin by first briefly establishing the definition of a **topological n -manifold** as a Hausdorff topological space M that is second countable (i.e., it has a countable basis) with the property that every point in M has an open neighborhood that is homeomorphic to either $\mathbb{R}^n$ or $\mathbb{H}^n$ where $\mathbb{H}^n$ is the *half space* defined as $\mathbb{H}^n = \{(x_1, \dots, x_n); x_n \geq 0\}$. The **boundary** of M , denoted as ∂M is the set of all points in M that have a neighborhood homeomorphic to an open set of the half space $\mathbb{H}^n$ taking the point to $\partial \mathbb{H}^n = \{(x_1, \dots, x_n); x_n = 0\}$. M will be said to have **no boundary** if $\partial M = \emptyset$, i.e., if every element of M is homeomorphic to $\mathbb{R}^n$. Also, we say that M is **closed** when it is compact with no boundary.

Now, with the objective of establishing the most important theorem of this section, the Poincaré duality, we state, without proof a series of lemmas that are useful in proving such a theorem, and also they may help us study homology and cohomology of manifolds.

Lemma A.1.5. *If U is an open subset of $\mathbb{R}^n$ then $H_i(U) = 0$ for every $i \geq n$.*

Lemma A.1.6. *If M is an n -manifold without boundary, then $H_i(M) = 0$ for $i > n$.*

From the previous lemma we get that the non trivial homology groups of a manifold occur only on dimensions less than or equal to the dimension of the manifold. We will proceed to conclude that if M is an n -dimensional manifold connected but not compact, then the n -th homology group is also trivial.

Lemma A.1.7. *Let U be an open set of $\mathbb{R}^n$ and $a \in H_n(\mathbb{R}^n, U)$. If for every $p \in \mathbb{R}^n - U$, the homomorphism induced by inclusion*

$$j_{p*}: H_n(\mathbb{R}^n, U) \rightarrow H_n(\mathbb{R}^n, \mathbb{R}^n - p)$$

has $J_{p}(a) = 0$ then $a = 0$.*

Theorem A.1.13. *If M is connected, non compact n -manifold without boundary, then $H_n(M) = 0$.*

Corolary A.1.2. *Let M be a closed connected n -manifold. If $z \in H_n(M)$ and $p \in M$ such that*

$$i_{p*}: H_n(M) \rightarrow H_n(M, M - p)$$

has $i_{p}(z) = 0$, then $z = 0$.*

Corolary A.1.3. *If M is a connected manifold without boundary, then either*

i. $H_n(M) = 0$ or

ii. $H_n(M) \simeq \mathbb{Z}$ and for every $p \in M$ the homomorphism $i_{p}: H_n(M) \rightarrow H_n(M, M - p)$ is an isomorphism.*

Consider first an n -manifold M , from calculus and differential geometry we have a brief notion of what an orientation must be, but now we want to be able to talk about orientations in the context of algebraic topology, and more specifically, in the context of manifolds. Consider a point p in M , we define a **local orientation** of M in the point p , a choice of generator for the infinite group $H_n(M, M - \{p\})$. Note that taking U an open set of M around p that is homeomorphic to $\mathbb{R}^n$, we get by using the excision theorem the following:

$$H_n(M, M - \{p\}) \simeq H_n(U, U - \{p\}) \simeq H_n(\mathbb{R}^n, \mathbb{R}^n - \{p\})$$

now, by the exact sequence of pairs, we have:

$$\cdots \rightarrow H_n(\mathbb{R}^n) \rightarrow H_n(\mathbb{R}^n, \mathbb{R}^n - \{p\}) \xrightarrow{\Delta} H_{n-1}(\mathbb{R}^n - \{p\}) \rightarrow H_{n-1}(\mathbb{R}^n) \rightarrow \cdots$$

now, since $\mathbb{R}^n$ is contractible and $\mathbb{R}^n - \{p\}$ is homeomorphic to S^{n-1} we get that Δ is an isomorphism, hence:

$$H_n(M, M - \{p\}) \simeq H_{n-1}(S^{n-1}) \simeq \mathbb{Z}.$$

So we have two choices of generators for $H_n(M, M - \{p\})$, i.e., two local orientations for M at the point p .

Now we extend this notion of local orientation to a global orientation, and we can do that by simply saying that a (global) **orientation** is a function assigning each point $x \in M$ to a local orientation $\mu_x \in H_n(M, M - \{p\})$ and that satisfy: for every $x \in M$ there exists U open contained in M that contains an open ball, B , around x and such that all the local orientations μ_y at the point y in this open ball are the images of the generator μ_x under the natural inclusion maps $H_n(M, M - B) \rightarrow H_n(M, M - \{y\})$.

In other words, an orientation for M is a choice of local orientation that is compatible with all local orientations around each point of M .

The notion of orientation on a manifold can be extended to the notion of **R-orientation** by changing the coefficients of the homology groups involved from $\mathbb{Z}$ to a commutative ring with unit R , thus we may say that an R -orientation for M is a choice of generator for $H_n(M, M - \{p\}; R)$ that is compatible with local orientations.

Theorem A.1.14. *Let M be a closed connected manifold, we have that*

i. If M is R orientable, then the map $H_n(M; R) \rightarrow H_n(M, M - \{p\}) \simeq R$ is an isomorphism for all $p \in M$;

- ii. If M is not R orientable, then the map $H_n(M; R) \rightarrow H_n(M, M - \{p\}) \simeq R$ is injective and its image is the set $R_2 = \{r \in R; 2r = 0\}$;
- iii. $H_i(M; R) = 0$ for all $i > n$.

The following proposition uses the previous result, and it is very helpful when trying to classify homology of manifolds, as we will see later:

Proposition A.1.4. *If M is a closed connected manifold of dimension n the torsion subgroup of H_{n-1} is either trivial if M is orientable or $\mathbb{Z}_2$ if M is non-orientable.*

Proof. Suppose M is orientable and $H_{n-1}(M)$ is not torsion-free, i.e., we can write $H_{n-1}(M; \mathbb{Z}) \simeq F \oplus \text{Tor}(H_{n-1}(M; \mathbb{Z}))$ for F a free abelian group, then, for some prime p , by using the universal coefficients theorem, we would get $H_n(M; \mathbb{Z}_p)$ would be larger than $\mathbb{Z}_p$, since:

$$\begin{aligned} H_n(M; \mathbb{Z}_p) &\simeq H_n(M; \mathbb{Z}) \otimes \mathbb{Z}_p \oplus \text{Tor}(H_{n-1}(M; \mathbb{Z}), \mathbb{Z}_p) \\ &\simeq \mathbb{Z}_p \oplus \text{Tor}(F, \mathbb{Z}_p) \oplus \text{Tor}(\mathbb{Z}_p, \mathbb{Z}_p) \\ &\simeq \mathbb{Z}_p \oplus \mathbb{Z}_p \end{aligned}$$

but, since M is supposed orientable, we should get $H_n(M; \mathbb{Z}_p) \simeq \mathbb{Z}_p$, hence $H_{n-1}(M; \mathbb{Z})$ has trivial torsion.

Now, if M is non orientable, we know that $H_n(M; \mathbb{Z}_m)$ is either $\mathbb{Z}_2$ or 0, note that if m is odd, since $H_n(M; R) \rightarrow R$ is injective with image $2R = \{r \in R; 2r = 0\}$, then the image of $H_n(M; \mathbb{Z}_m) \rightarrow \mathbb{Z}_m$ must be zero, because no element in $\mathbb{Z}_m$ has order 2. Now, if m is even, then there will be two elements in $\mathbb{Z}_m$ with order 2, namely $\bar{0}$ and $\frac{\bar{m}}{2}$, hence $H_n(M; \mathbb{Z}_m) \simeq \mathbb{Z}_2$. Suppose then, that H_{n-1} has some torsion other than $\mathbb{Z}_2$, then there must exist some $p > 2$ such that by the universal coefficients theorem $H_n(M; \mathbb{Z}_p) \simeq \text{Tor}(H_{n-1}(M), \mathbb{Z}_p)$ is greater than $\mathbb{Z}_k$ (similarly to the orientable case), which is a contradiction, because $H_n(M, \mathbb{Z}_m)$ must be either $\mathbb{Z}_2$ or 0. Hence, $H_{n-1}(M)$ has torsion $\mathbb{Z}_2$. $\square$

Proposition A.1.5. *Let M be a closed manifold of odd dimension, then the Euler characteristic of M is zero, i.e., $\chi(M) = 0$.*

Proof. From Poincaré duality, we have that if M is orientable, then $H_i(M; \mathbb{Z}) \simeq H^{n-i}(M, \mathbb{Z})$, where n is the dimension of M , hence $\text{rank}(H_i(M; \mathbb{Z})) = \text{rank}(H^{n-i}(M, \mathbb{Z}))$, and from the universal coefficients theorem we get that

$$\text{rank}(H^{n-i}(M; \mathbb{Z})) = \text{rank}(H_{n-i}(M; \mathbb{Z})).$$

Thus, if we assume that n is odd, all the terms of:

$$\chi(M) = \sum_{i=0}^n (-1)^i \text{rank}(H_i(M; \mathbb{Z}))$$

cancel out, hence $\chi(M) = 0$.

If M is non-orientable, note that $\text{rank}(H_n(M, \mathbb{Z})) = \dim_{\mathbb{Z}_2}(H_i(M; \mathbb{Z}_2))$, hence, we may apply the same argument to get:

$$\sum_{i=0}^n (-1)^i \dim_{\mathbb{Z}_2}(H_i(M; \mathbb{Z}_2)) = 0.$$

What we have to check now, is that the above sum equals the Euler characteristic of M . In fact, from the universal coefficient theorem for cohomology we get:

$$H^i(M; \mathbb{Z}_2) \simeq \text{Hom}(H_i(M; \mathbb{Z}_2), \mathbb{Z}_2) \oplus \text{Ext}(H_{i-1}(M; \mathbb{Z}_2), \mathbb{Z}_2)$$

we get $H^i(M; \mathbb{Z}_2) \simeq H_i(M; \mathbb{Z}_2)$ since $\mathbb{Z}_2$ is a field and $H_i(M; \mathbb{Z}_2)$ is finitely generated, hence isomorphic to its dual (view as a $\mathbb{Z}_2$ -vector space). Now each $\mathbb{Z}$ summand of $H_i(M; \mathbb{Z})$ gives a $\mathbb{Z}_2$ summand of $H^i(M; \mathbb{Z}_2)$, each $\mathbb{Z}_m$ summand of $H_i(M; \mathbb{Z})$ gives a $\mathbb{Z}_2$ summand of $H^i(M; \mathbb{Z}_2)$ and $H^{i+1}(M; \mathbb{Z}_2)$ when m is even, and zero summands when m is odd, hence the contributions of these $\mathbb{Z}_2$ summands given by m even will cancel out on the sum since consecutive terms have opposite signs. Hence, we have the equality

$$\chi(M) = \sum_{i=0}^n (-1)^i \dim_{\mathbb{Z}_2}(H_i(M; \mathbb{Z}_2)).$$

□

Finally we are able to establish the Poincaré Duality:

Theorem A.1.15 (Poincaré Duality). *For M a compact connected orientable n -manifold without boundary, the homomorphism:*

$$D: H^k(M; G) \rightarrow H_{n-k}(M; G)$$

given by $D(x) = x \frown z$ is an isomorphism for each k . Here z is the fundamental class associated to the orientation chosen for M .

If M is an n -manifold with boundary, there is a version of Poincaré duality, that is given in the following theorem:

Theorem A.1.16 (Poincaré-Lefschetz Duality). *Let M be an n -manifold with boundary ∂M and R a commutative ring with unit. Then*

$$H^q(M; R) \simeq H_{n-q}(M, \partial M; R) \quad H^q(M, \partial M; R) \simeq H_{n-q}(M; R).$$

A.2 Low Dimensional Manifolds

We now verify some properties and results obtained from the study of homology on manifolds. The results involving cohomology end up being the same by usage of Poincaré duality. From now on, unless explicit, all manifolds will be considered connected, since non connected manifolds are nothing but disjoint union of connected ones, hence, its homologies are direct sums of homologies of each connected component.

2-Manifolds

Consider first an 2-manifold without boundary S , also called surface. Recall that from the classification theorem of surfaces we have that if S is orientable then S is homeomorphic to S^2 , $\mathbb{T}^2$ or Σ_g the connected sum of g tori, we shall also denote Σ_g by $\#g\mathbb{T}^2$. Hence, we have to determine the homology groups of each of these surfaces. In fact, we already have the homology groups for S^2 and $\mathbb{T}^2$, to determine $H_*(\Sigma_g)$ we shall use the following result:

Proposition A.2.1. *Let M and N be two connected n -manifolds. If M or N is orientable then $H_i(M \# N) = H_i(M) \oplus H_i(N)$ for $0 < i < n$. If they are both non-orientable, then $H_i(M \# N)$ is obtained from $H_i(M) \oplus H_i(N)$ by replacing one of the $\mathbb{Z}_2$ summands by a $\mathbb{Z}$ summand.*

Hence we get the following

$$H_1(\Sigma_g) = H_1(\Sigma_{g-1}) \oplus H_1(\mathbb{T}^2)$$

and supposing $H_1(\Sigma_{g-1}) \simeq \mathbb{Z}^{2(g-1)}$ (we can do that, since we have already computed the homology for $g = 1$ and for $g = 2$, then we are proceeding by induction) we get

$$H_1(\Sigma_g) = \mathbb{Z}^{2g-2} \oplus \mathbb{Z} \oplus \mathbb{Z} = \mathbb{Z}^{2g}$$

and we get:

$$H_n(\Sigma_g) \simeq \begin{cases} \mathbb{Z}, & \text{if } n = 0, 2 \\ \mathbb{Z}^{2g}, & \text{if } n = 1 \\ 0, & \text{otherwise} \end{cases}.$$

And we have classified the homology for orientable surfaces without boundary.

We now consider the case where S is non orientable, then by the same classification theorem for surfaces, S is homeomorphic to either $\mathbb{R}P(2)$ or $\#k\mathbb{R}P(2)$. We already know that

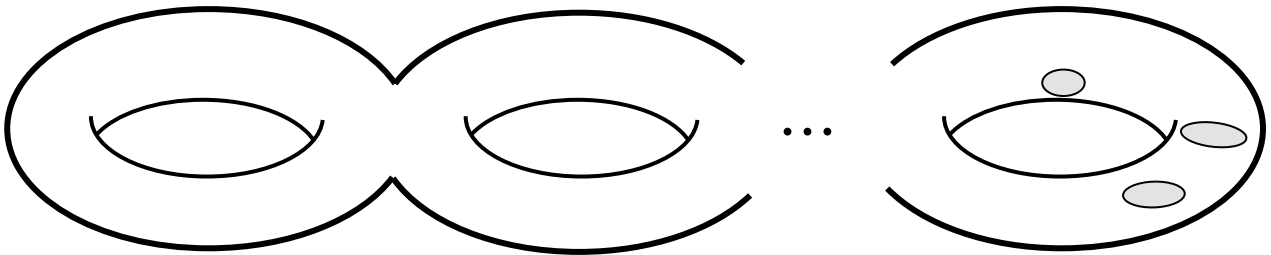
$$H_n(\mathbb{R}P(2)) \simeq \begin{cases} \mathbb{Z}, & \text{if } n = 0 \\ \mathbb{Z}_2, & \text{if } n = 1 \\ 0, & \text{otherwise} \end{cases}$$

thus, by using the Proposition A.2.1, we get:

$$H_n(\#k\mathbb{R}P(2)) \simeq \begin{cases} \mathbb{Z}, & \text{if } n = 0 \\ \mathbb{Z}^{k-1} \oplus \mathbb{Z}_2, & \text{if } n = 1 \\ 0, & \text{otherwise} \end{cases}.$$

Next step is to deal with orientable manifolds with boundary. To generalize the previous classification theorem but now with boundary, we say that an orientable surface with boundary is homeomorphic to either a disk (which is contractible, therefore not very interesting) or to $\Sigma_{g,b}$ where g is the genus of the surface (i.e., the amount of tori we are gluing together) and b is the number of open disks removed from the interior of Σ_g .

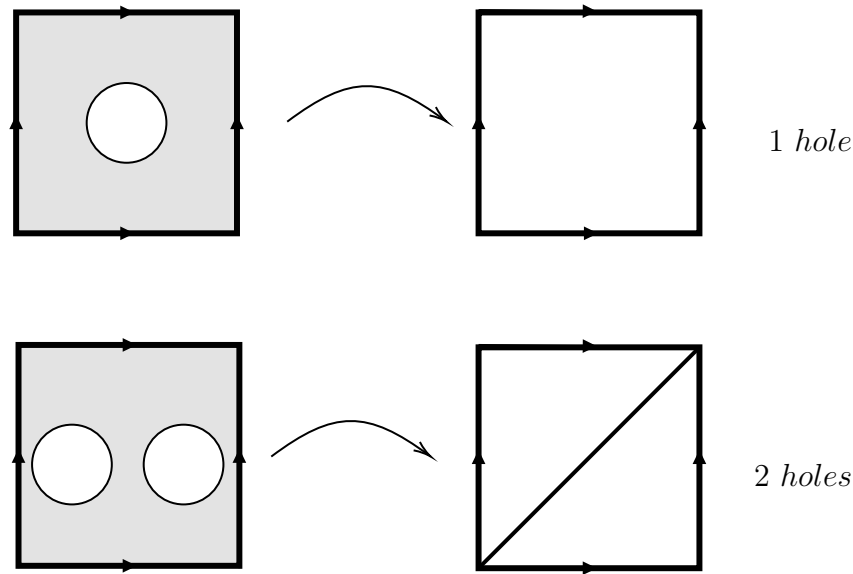
The first thing we have to note is that all the open disks removed can be homeomorphically deformed to lie all in the same "torus component" of Σ_g , as in the picture below.



Now, we compute first $\Sigma_{1,1}$, then use a similar argument to get $\Sigma_{1,b}$ and complete the calculation by using Proposition A.2.1. It is obvious that $\Sigma_{1,1}$ is just a torus with an open disk removed, thus it deformation retracts to $S^1 \vee S^1$, hence

$$H_1(\Sigma_{1,1}) \simeq \mathbb{Z} \oplus \mathbb{Z}.$$

Now, consider $\Sigma_{1,2}$, this can be seen as the planar representation of the torus, but with two open disks removed from its interior, we can deform these holes to get a similar picture, but now it is not "solid" and it has a segment going from the left vertex the bottom segment to the right vertex of the top segment. When we glue this picture as indicated by the identified edges we obtain the wedge product of 3 circles, as in the picture:



hence $H_1(\Sigma_{1,2}) \simeq \mathbb{Z} \oplus \mathbb{Z} \oplus \mathbb{Z}$.

Assume now that we have b holes in the square that defines the torus. We can deform these holes in a similar fashion of that done for two holes, but now we get $b - 1$ new segments instead of one, closing this picture yields us $\bigvee_{i=1}^{b-1} S^1$, hence

$$H_1(\Sigma_{1,b}) = H_1\left(\bigvee_{i=1}^{b-1} S^1\right) \simeq \mathbb{Z}^{b-1}.$$

To conclude, we use Proposition A.2.1 to get

$$H_n(\Sigma_{g,b}) \simeq \begin{cases} \mathbb{Z}, & \text{if } n = 0 \\ \mathbb{Z}^{2g+b-1}, & \text{if } n = 1 \\ 0, & \text{otherwise} \end{cases}.$$

3-Manifolds

Going up a dimension we now discuss a little bit about 3-manifolds. We shall first classify the closed manifolds without boundary in a general way and from there we discuss those which have trivial fundamental group.

Note that the Hurevich's theorem is established for path connected topological spaces, in particular, connected manifolds are also path connected, hence the first homology

group of 3-manifolds is already determined by this theorem. Of course, the 0-th and third homologies are also determined as always. We shall then study the second homology group.

Proposition A.2.2. *If $H_1(M) \simeq \mathbb{Z}^r \oplus \text{Tor}(H_1(M))$ then $H_2(M) \simeq \mathbb{Z}^r$ if M is orientable and $H_2(M) \simeq \mathbb{Z}^{r-1} \oplus \mathbb{Z}_2$ if M is non-orientable.*

Proof. Recall that $\text{rank}(H_n(M))$ is the number of generators of the torsion-free part of the decomposition of $H_n(M)$. Then, from what we know about the homology groups of M we get:

- $H_0(M) \simeq \mathbb{Z} \Rightarrow \text{rank}(H_0(M)) = 1$;
- $H_1(M) \simeq \mathbb{Z}^r \oplus \text{Tor}(H_1(M))$ from hypothesis, hence $\text{rank}(H_1(M)) = r$;
- $H_3(M) \simeq \begin{cases} \mathbb{Z}, & \text{if } M \text{ is orientable} \\ 0, & \text{if } M \text{ is non-orientable} \end{cases} \Rightarrow \text{rank}(H_3(M)) = \begin{cases} 1, & \text{if } M \text{ is orientable} \\ 0, & \text{if } M \text{ is non-orientable} \end{cases}$

From that and from the fact that the Euler characteristic of a closed manifold of dimension odd is always zero, we get the following cases:

$$\text{rank}(H_2(M; \mathbb{Z})) = \begin{cases} r, & \text{if } M \text{ is orientable} \\ r - 1, & \text{if } M \text{ is non-orientable} \end{cases}.$$

Now we just need to deal with the torsion part of each case. But from Proposition A.1.4 we have that $H_2(M)$ is torsion free when it is orientable, and has torsion $\mathbb{Z}_2$ when it is non-orientable, this concludes the result. $\square$

From that, we get the following about the homology of closed and orientable 3-manifolds:

$$H_n(M; \mathbb{Z}) \simeq \begin{cases} \mathbb{Z}, & \text{if } n = 0, 3 \\ \pi_1(M)_{Ab}, & \text{if } n = 1 \\ \frac{\pi_1(M)_{Ab}}{\text{Tor}(\pi_1(M)_{Ab})}, & \text{if } n = 2 \\ 0, & \text{otherwise} \end{cases}.$$

From that we can already classify the simply connected 3-manifolds, i.e., those which have trivial fundamental groups, we get

$$H_n(M; \mathbb{Z}) \simeq \begin{cases} \mathbb{Z}, & \text{if } n = 0, 3 \\ 0, & \text{otherwise} \end{cases}.$$

Concerning 3-manifolds we also have some interesting results by assuming something about its first homology group. For example if M is a 3-manifold with boundary such that $H_1(M)$ is finite we can use the universal coefficients theorem to get:

$$H_1(M; \mathbb{Q}) \simeq H_1(M; \mathbb{Z}) \otimes \mathbb{Q} \oplus \text{Tor}(H_0(M; \mathbb{Z}), \mathbb{Q})$$

and since the tensor product of any finite group with $\mathbb{Q}$ is zero, and $H_0(M)$ is free abelian, we get $H_1(M; \mathbb{Q}) = 0 \Rightarrow H_2(M; \mathbb{Q}) \simeq \frac{0}{\text{Tor}(0)} = 0$.

Now, considering the pair $(M, \partial M)$ we get the following long exact sequence of pairs:

$$\cdots \rightarrow \underline{H_2(M; \mathbb{Q})} \xrightarrow{0} H_1(\partial M; \mathbb{Q}) \rightarrow \underline{H_1(M; \mathbb{Q})} \xrightarrow{0} \cdots$$

therefore $H_1(\partial M; \mathbb{Q}) \simeq 0$, hence, the boundary of M must be a disjoint union of copies of S^2 .

Also, using this previous result, one gets the contrapositive, i.e., if M has no copies of S^2 in its boundary, then its first homology must be infinite.

Lens Spaces

The lens spaces are going to be used later to construct spaces with arbitrary torsion groups that will be useful in studying for example the *linking form* that are bilinear forms defined on torsion subgroups of homology groups. Take S^3 the usual 3-sphere and we shall view it as a subspace of $\mathbb{C}^2$, i.e., $S^3 = \{(z_1, z_2) \in \mathbb{C}^2; |z_1| + |z_2| = 1\}$. Now, let $\mathbb{Z}_m$ be the usual finite group of order m and define a relation

$$(z_1, z_2) \sim (z'_1, z'_2) \Leftrightarrow (z'_1, z'_2) = (\zeta_m z_1, \zeta_m^n z_2)$$

where ζ_m is the m -th root of unity where m is coprime with n . Then we define the lens space $L(m; n)$ as

$$L(m; n) = \frac{S^3}{\mathbb{Z}_m} = \frac{S^3 = \{(z_1, z_2) \in \mathbb{C}^2; |z_1| + |z_2| = 1\}}{\sim}.$$

If we write the complex number in their polar form, we can also write:

$$L(m; n) = \frac{\{(r_1 e^{2\pi i \theta_1}, r_2 e^{2\pi i \theta_2}; r_1^2 + r_2^2 = 1\}}{(\theta_1, \theta_2) \sim \left(\theta_1 + \frac{1}{m}, \theta_2 + \frac{n}{m}\right)}.$$

If $\rho: S^3 \rightarrow S^3$ is given by $(z_1, z_2) \mapsto (\zeta_m z_1, \zeta_m^n z_2)$, then we have an action of the group $\mathbb{Z}_m$ on S^3 , that takes $\bar{k} \in \mathbb{Z}_m$ into ρ^k , then $L(m; n)$ can be viewed as the quotient of S^3 by an action of the group $\mathbb{Z}_m$.

Important features of the lens spaces are that $L(m; n)$ is homotopy equivalent to $L(m', n')$ if, and only if, $m = m'$ and $n' \equiv n^{\pm 1} k^2 \pmod{m}$ for some $k \in \mathbb{N}$. Also, $L(m; n)$ is homeomorphic to $L(m', n')$ if, and only if $m = m'$ and $n' \equiv \pm n^{\pm 1} \pmod{m}$.

Since S^3 is the universal cover for $L(m; n)$ we get that its fundamental group $\pi_1(L(m; n)) \simeq \mathbb{Z}_m$, and since this group is abelian, from Hurevich's theorem ($L(m; n)$ is orientable) we get the same for its first homology group, i.e., $H_1(L(m; n); \mathbb{Z}) \simeq \mathbb{Z}_m$.

With a quick look at these results one can easily see that both homology groups and the fundamental group does not depend on the choice of the action of $\mathbb{Z}_m$ on S^3 , since neither depends on n , although for two lens spaces to be homotopic or homeomorphic, the choice of the action must be taken into consideration. This shows that homology and fundamental groups are not enough to determine certain characteristics of the lens spaces.

4-Manifolds

We start by looking at closed orientable simply connected manifolds without boundary M . Observe that the hypothesis of being simply connected implies that (as in the case

of 3-manifolds) the fundamental group of M is trivial, hence its first homology group is also trivial by Hurevich's theorem. Now, of course $H_0(M) \simeq H_4(M) \simeq \mathbb{Z}$ so these cases are not very interesting.

If we use the Poincaré duality here, we get $H_3(M) \simeq H^1(M)$ and by the universal coefficients theorem for cohomology, we get:

$$H_3(M) \simeq \text{Hom}(H_1(M), \mathbb{Z}) \oplus \text{Ext}(H_0(M), \mathbb{Z})$$

and since $H_1(M) \simeq 0$ and $\mathbb{Z}$ is free abelian, we get $H_3(M) = 0$ as well.

The most interesting case here is for $H_2(M) \simeq H^2(M)$. Note that using again the universal coefficients theorem for cohomology, we may write:

$$H_2(M) \simeq \text{Hom}(H_2(M), \mathbb{Z}) \oplus \text{Ext}(H_1(M), \mathbb{Z})$$

again, $\mathbb{Z}$ is free abelian, and $H_1(M) = 0$ hence $\text{Ext}(0, \mathbb{Z}) = 0$. Now, note that H_2 is abelian therefore $\text{Hom}(H_2(M), \mathbb{Z})$ is a torsion-free group, i.e., $H_2(M)$ is free. Hence we get the following classification for closed simply connected orientable 4-manifolds:

$$H_n(M) = \begin{cases} \mathbb{Z}, & \text{if } n = 0, 4 \\ \mathbb{Z}^k & \text{for some } k \in \mathbb{Z}_+, \text{ if } n = 2 \\ 0 & \text{otherwise} \end{cases}.$$

An important tool in the study of 4-manifolds is the intersection pairing, that is induced by the cup product:

$$\begin{aligned} \smile: H^p(M) \otimes H^q(M) &\rightarrow H^{p+q}(M) \\ ([\alpha], [\beta]) &\mapsto [\alpha \smile \beta] \end{aligned}$$

Note that in the middle homology of an orientable 4-manifold M , $H_2(M)$ this pairing is specifically interesting because it sends pairs of classes of 2-cochains into the class of 4-cochains, that can be seen as an integer since $H_4(M) \simeq \mathbb{Z}$. Geometrically, this pairing gives the number of intersection between two surfaces embedded in M that generate $H_2(M)$. To see this, we use Poincaré duality to identify each cohomology group to its homology group, which in turn has elements that can be seen as such embedded surfaces.

We formally define the intersection pairing as follows:

$$\begin{aligned} Q_M: H^2(M) \times H^2(M) &\rightarrow H^4(M) \simeq \mathbb{Z} \\ (\alpha, \beta) &\mapsto \langle \alpha \smile \beta, [M] \rangle \end{aligned}$$

where $[M]$ is the fundamental class of M , i.e., the class that is related to the choice of orientation on M . As noted previously, if α is the class of an embedded surface A and β the class of an embedded surface B , then $Q_M(\alpha, \beta) = [A \cap B]$. We note that by choosing a basis $\{\alpha_i\}$ for $H_2(M)$ we may associate to Q_M a matrix, since Q_M is a bilinear form (the cup product is bilinear).

To see an example, consider $M = S^2 \times S^2$. We first compute the homology of M by using Künneth formula:

$$\begin{aligned} H_2(S^2 \times S^2) &\simeq \bigoplus_{p+q=2} H_p(S^2) \otimes H_q(S^2) \oplus \bigoplus_{r+s=1} \text{Tor}(H_r(S^2), H_s(S^2)) \\ &\simeq H_0(S^2) \otimes H_2(S^2) \oplus H_1(S^2) \otimes H_1(S^2) \oplus H_2(S^2) \otimes H_0(S^2) \oplus \\ &\quad \oplus \text{Tor}(H_0(S^2), H_1(S^2)) \oplus \text{Tor}(H_1(S^2), H_0(S^2)) \\ &\simeq \mathbb{Z} \otimes \mathbb{Z} \oplus 0 \times 0 \oplus \mathbb{Z} \otimes \mathbb{Z} \oplus \text{Tor}(\mathbb{Z}, 0) \oplus \text{Tor}(\mathbb{Z}, 0) \\ &\simeq \mathbb{Z} \oplus \mathbb{Z} \end{aligned}$$

Now, we take $\alpha = [S^2 \times \{q\}]$ and $\beta = [\{p\} \times S^2]$ note that $\{\alpha, \beta\}$ is a basis for $H_2(M) \simeq \mathbb{Z} \oplus \mathbb{Z}$. And by what we have said before, we get:

$$Q_{S^2 \times S^2}(\alpha, \beta) = [(S^2 \times \{q\}) \cap (\{p\} \times S^2)] = [(p, q)] = \mu$$

where μ is the class of a point in $H_4(M)$ which is identified with 1 in $\mathbb{Z}$. Analogously we can obtain $Q_{S^2 \times S^2}(\beta, \alpha) = \mu$. Now, to compute the pairing $Q_{S^2 \times S^2}(\alpha, \alpha)$ we need to choose two different representatives of the class α , for example $S^2 \times \{p\}$ and $S^2 \times \{p'\}$, of course $(S^2 \times \{p\}) \cap (S^2 \times \{p'\}) = \emptyset$, which is identified with $0 \in \mathbb{Z}$, hence $Q_{S^2 \times S^2}(\alpha, \alpha) = 0$, same thing happens for $Q_{S^2 \times S^2}(\beta, \beta) = 0$.

Now, we can organize these results in a matrix, since Q_M is a bilinear form, it has an $n \times n$ matrix associated with it, where n is the dimension of $H_2(M)$. In the previous example, we get the following associated matrix, that will be denoted as $[Q_M]$:

$$[Q_{S^2 \times S^2}] = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}.$$

If we denote by b_+ the number of positive eigenvalues of Q_M and by b_- the number of negative eigenvalues of Q_M , then we can define the **signature** of M as the signature of the intersection form of M :

$$\sigma(M) = b_+ - b_-.$$

Hence, if we go back to the example, we first compute its eigenvalues:

$$\det \begin{pmatrix} -\lambda & 1 \\ 1 & -\lambda \end{pmatrix} = 0 \Leftrightarrow \lambda^2 - 1 = 0 \Leftrightarrow \lambda = \pm 1$$

hence $b_+ = b_- = 1$ and therefore the signature of $S^2 \times S^2$ is $\sigma(S^2 \times S^2) = 1 - 1 = 0$.

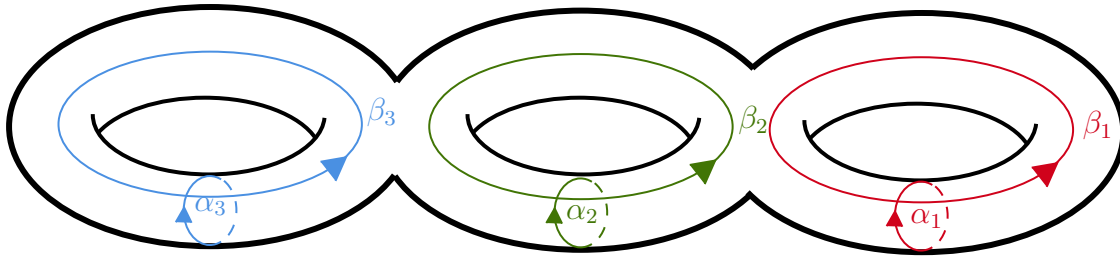
The signature of manifolds satisfies interesting properties:

- If $\partial M = \varepsilon$ and $M = M_1 \cup_{\partial} M_2$, then $\sigma(M) = \sigma(M_1) + \sigma(M_2)$;
- If we denote by $\overline{M}$ the manifold M with opposite orientation, then $\sigma(\overline{M}) = -\sigma(M)$;
- If $M = \partial N$, i.e., M is the boundary of a 5-manifold N compact and oriented, then $\sigma(M) = 0$;
- $\sigma(M \times N) = \sigma(M)\sigma(N)$.

From the third property we conclude that if a 4-manifold has non zero signature, than it cannot bound a compact oriented 5-manifold.

We do not need to restrict ourselves to 4-manifolds when studying intersection forms. In fact, we may also study intersection forms on surfaces, which have a more geometrical appealing, for example, the middle homology of a surface is $H_1(S)$, whose elements may be seen as closed curves in the surface, hence the intersection form on surfaces, gives us the number of intersection between two closed curves on the basis of $H_1(S)$. We shall now do an example to clarify things.

Take, for example, $S = \Sigma_3$. We have that $H_1(S)$ have 6 generators, i.e., it has a basis $\{\alpha_1, \alpha_2, \alpha_3, \beta_1, \beta_2, \beta_3\}$ that are the the class of the loops of the type $\alpha_i = S^1 \times \{q_i\}$ and $\beta_i = \{p_i\} \times S_1$ for each $i = 1, 2, 3$. As in the picture below.



We will not write down all the computations, but we can see from the picture that the α_i 's do not intersect each other, as well as the β_i 's do not intersect each other, also α_i only intersects β_i and only at one point, hence $Q_{\Sigma_3}(\alpha_i, \beta_i) = Q_{\Sigma_3}(\beta_i, \alpha_i) = \mu$ for each $i = 1, 2, 3$. To see that each α_i (respec. β_i) do not intersect itself, we take two different generators for the same 1-cycle, and note that they do not intersect in the same fashion as we did previously for $S^2 \times S^2$. Hence, we get the following matrix for the form Q_{Σ_3} :

$$[Q_{\Sigma_3}] = \begin{pmatrix} 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \end{pmatrix}$$

which has signature $\sigma(\Sigma_3) = 0$.