



Universidad de San Carlos de Guatemala
Escuela de Ciencias Físicas y Matemáticas
Departamento de Matemática

ESTUDIO TEÓRICO SOBRE INVARIANTES DE NUDOS Y SUS CONEXIONES CON FÍSICA DE PARTÍCULAS

Byron Abel Raul Hernández Pacay

Asesorado por Lic. Billy Othmaro Quevedo Carranza

Guatemala, mayo de 2025

UNIVERSIDAD DE SAN CARLOS DE GUATEMALA



ESCUELA DE CIENCIAS FÍSICAS Y MATEMÁTICAS

**ESTUDIO TEÓRICO SOBRE INVARIANTES DE
NUDOS Y SUS CONEXIONES CON FÍSICA DE
PARTÍCULAS**

TRABAJO DE GRADUACIÓN
PRESENTADO A LA JEFATURA DEL
DEPARTAMENTO DE MATEMÁTICA
POR

BYRON ABEL RAUL HERNÁNDEZ PACAY
ASESORADO POR LIC. BILLY OTHMARO QUEVEDO CARRANZA

AL CONFERÍRSELE EL TÍTULO DE
LICENCIADO EN MATEMÁTICA APLICADA

GUATEMALA, MAYO DE 2025

UNIVERSIDAD DE SAN CARLOS DE GUATEMALA
ESCUELA DE CIENCIAS FÍSICAS Y MATEMÁTICAS



CONSEJO DIRECTIVO INTERINO

Director	M.Sc. Jorge Marcelo Ixquiac Cabrera
Representante Docente	Arqta. Ana Verónica Carrera Vela
Representante Docente	M.A. Pedro Peláez Reyes
Representante de Egresados	Lic. Urías Amitaí Guzmán García
Representante de Estudiantes	Elvis Enrique Ramírez Mérida
Representante de Estudiantes	Oscar Eduardo García Orantes
Secretario	M.Sc. Freddy Estuardo Rodríguez Quezada

TRIBUNAL QUE PRACTICÓ EL EXAMEN GENERAL PRIVADO

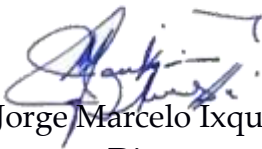
Director	M.Sc. Jorge Marcelo Ixquiac Cabrera
Examinador	M.Sc. Edgar Damián Ochoa Hernández
Examinador	M.Sc. Frank Jorge Frizstche Barrios
Examinador	Lic. José Daniel Romero
Secretario	M.Sc. Freddy Estuardo Rodríguez Quezada

Ref. D.DTG. M001-2025
Guatemala 13 de mayo de 2025

El Director de la Escuela de Ciencias Físicas y Matemáticas de la Universidad de San Carlos de Guatemala, luego de conocer la aprobación por parte del jefe de la Licenciatura en Matemática Aplicada, al trabajo de graduación titulado: "ESTUDIO TEÓRICO SOBRE INVARIANTES DE NUDOS Y SUS CONEXIONES CON FÍSICA DE PARTÍCULAS", presentado por el estudiante universitario Byron Abel Raul Hernández Pacay, autoriza la impresión del mismo.

IMPRÍMASE.

"ID Y ENSEÑAD A TODOS"



M.Sc. Jorge Marcelo Ixquiac Cabrera
Director

AGRADECIMIENTOS

A Dios

Por darme sabiduría, estar conmigo en mi carrera y librarme de tantas situaciones en estos años.

A mi padre

Por confiar en mis decisiones, apoyarme económicamente y enseñarme a resolver problemas.

A mi madre

Por estar para mí en mis momentos más difíciles, llevarme café en incontables noches de desvelo, sostener a la familia con sus oraciones y mostrarme el valor de la disciplina.

A mi asesor

Por dedicarme tiempo y tomar el reto de revisar detalladamente este trabajo.

A mis mentoras

Mónica Cabria e Isabela Recio Hernández, quienes me ayudaron a dar mis primeros pasos en topología y compartieron su conocimiento sin nada de egoísmo conmigo.

A mi grupo de amigos

Los dos Diegos, William, Tavo, Jorge, Ana, David, Darián, Gabriel, Isabel, Jonathan, Merie, Rodrigo, Marielos, Christian, Dylan y Javier. Quienes me mostraron el valor de una amistad sincera y genuina.

A Audrick Veliz

Quien durante mis años en la universidad me motivó a ser un mejor matemático y compartimos frustraciones, triunfos así como muchas historias.

A la gente valiente

Quienes defendieron la autonomía de la USAC durante estos años.

DEDICATORIA

A mi tío Raúl Pacay, el primer hombre de la familia Pacay en llegar a la universidad y quien en vida fue un ejemplo de como se debe amar a la familia.

ÍNDICE GENERAL

ÍNDICE DE FIGURAS	IV
LISTA DE SÍMBOLOS	V
OBJETIVOS	VII
INTRODUCCIÓN	IX
1. NUDOS Y ENLACES	1
1.1. Teoría de nudos clásica	1
1.2. Teoría de nudos en la 3-esfera	12
1.3. Teoría general de nudos	17
2. TEOREMA DE SEIFERT-VAN KAMPEN	19
2.1. Equivalencia Homotópica	19
2.1.1. Homotopía	19
2.1.2. Tipo de homotopía	21
2.2. Grupo fundamental	26
2.2.1. Grupoide fundamental	26
2.2.2. Definición y propiedades de π_1	30
2.2.3. Espacios simplemente conexos	35
2.3. Presentaciones de grupos	37
2.4. Teorema de Van-Kampen	43
3. INVARIANTES ALGEBRAICOS	55
3.1. Grupo de un nudo	55
3.2. Estados y brackets	61
3.3. Polinomios de Jones y de Alexander-Conway	66
3.4. Homología de Khovanov	67

4. CONEXIONES CON FÍSICA	71
4.1. Funciones de Partición	71
4.2. Tensores abstractos y la ecuación de Yang-Baxter	74
4.3. Teorema de la estadística del espín	79
CONCLUSIONES	81
RECOMENDACIONES	83
BIBLIOGRAFÍA	85
A. TOPOLOGÍA GENERAL	89
B. HOMOLOGÍA	97
B.1. Definiciones generales	97
B.2. Homología singular	99

ÍNDICE DE FIGURAS

1.1. Cruces en diagramas para un nudo.	2
1.2. Dos representaciones del nudo trébol.	2
1.3. Cruces no permitidos de nudos en posición regular.	3
1.6. Idea de la composición de nudos.	7
1.7. Nudo trébol orientado	9
1.8. Operación con el nudo trivial.	9
1.9. Dos nudos distintos con mismos factores.	10
1.10. Triqueta con un círculo	10
1.11. Δ -Movimiento.	11
1.12. Movimientos de Reidemeister	12
1.13. Simplejos en $\mathbb{R}^3$	13
1.14. Nudo salvaje	15
1.15. Movie moves	17
2.1. Homotopía relativa a un conjunto A	20
2.2. Equivalencia de la banda de Möbius	22
2.3. Curvas cerradas en el cilindro.	24
2.4. Homomorfismos inducidos por aplicaciones homotópicas.	33
2.5. Deformación retráctil de un conjunto en forma de estrella.	35
3.1. Aplicación de Van-Kampen en $\mathbb{S}^2 \vee \mathbb{S}^1$	57
3.2. Generadores orientados en una vecindad de un cruce.	57
3.3. Lazos homotópicos al punto base en función de los generadores.	58
3.4. Túnel rectangular $\mathcal{N}$ usando un nudo como base.	59
3.5. Intersección para el uso del teorema de Van-Kampen.	59
3.6. Diagrama del enlace de Hopf orientado.	60
3.7. Descomposición de un cruce	61
3.8. Descomposición e identificación de un cruce.	61
3.9. Descomposición del nudo trébol por Kauffman	62

3.10. Cruces con orientación.	65
4.1. Procedimientos para el cálculo de un polinomio dicromático.	74
4.2. Tensor abstracto $T_{a_1 \dots a_m}^{b_1 \dots b_n}$	75
4.3. Diagrama como una matriz.	75
4.4. Operaciones de Matrices.	76
4.5. Cruces identificados.	76
4.6. Trébol orientado con cruces identificados.	76
4.7. Movimiento de Reidemeister 3 para un enlace orientado en el que todos sus cruces son positivos y están identificados.	77
4.8. Tensores abstractos utilizados.	78
4.9. Movimiento $\mathcal{R}_1$ aplicado a un cruce de una superficie.	79
4.10. Expresando cómo deshacer un cruce con las propiedades del bracket de Kauffman.	80

LISTA DE SÍMBOLOS

Símbolo	Significado
$:=$	es definido por
$\cong$	es isomorfo a
$\cong_{PL}$	PL-homeomorfismo
$\emptyset$	conjunto vacío
$E \setminus F$	diferencia entre E y F
$\mathbb{S}^n$	esfera n -dimensional
3_1	nudo trébol
8_{17}	nudo con 8 cruces número 17
a_b	nudo con a cruces número b
$C(X, Y)$	conjunto de funciones continuas de X a Y
$\simeq$	relación de homotopía
$\simeq_A$	relación de homotopía estacionaria
$\equiv$	equivalencia homotópica
$\Pi(X)$	grupoide fundamental
$\pi_1(X, p)$	grupo fundamental de X con base p
$\bigstar_{\alpha \in \Lambda} G_\alpha$	producto libre de grupos
$\text{Ncl}_G(H)$	clausura normal de H en G
$\langle L \rangle$	bracket de Kauffman
$[L]$	polinomio bracket no reducido
$\{L\}$	bracket especial
$T_{a_1 \dots a_m}^{b_1 \dots b_n}$	tensor abstracto
Grp	categoría de los grupos
Ring	categoría de los anillos
Ab	categoría de los grupos abelianos

OBJETIVOS

General

Describir las conexiones que existen entre la teoría de nudos y la física.

Específicos

1. Desarrollar los conceptos fundamentales de nudos.
2. Establecer las nociones de topología algebraica necesarias para el cálculo del grupo de un nudo a través del teorema de Seifert-Van Kampen.
3. Presentar algunos invariantes algebraicos de nudos bajo transformaciones de estos.
4. Exponer las conexiones que existen entre física de partículas y mecánica estadística con la teoría de nudos.

INTRODUCCIÓN

La teoría de nudos es un área relativamente joven de las matemáticas y relacionada desde sus inicios con fenómenos físicos. Sus primeros estudios datan aproximadamente de 1880, cuando William Thomson, mejor conocido como Lord Kelvin, consideró la hipótesis de que los átomos eran nudos en una sustancia llamada éter [1]. La meta principal de la teoría de nudos es determinar si dos nudos son equivalentes. Aunque esta meta parezca inocente, ha llevado a los matemáticos a buscar herramientas más sofisticadas para responder a las preguntas que conlleva. Algunas de estas herramientas provienen de la topología algebraica, como el grupo fundamental del complemento de un nudo a través de la presentación de Wirtinger. Asimismo, el estudio de los nudos ha dado como resultado la creación de algunas herramientas algebraicas como el polinomio bracket y el polinomio de Jones. Parte de estas construcciones son las que usaremos para presentar conexiones entre nudos y física, siendo la mayoría de estas expuestas originalmente por Kauffman en [18].

En el primer capítulo de este material se presentan conceptos generales dentro de la teoría de nudos con algunos detalles importantes sobre las definiciones, y se hace mención de un enfoque más moderno, luego en el capítulo 2; se hace una revisión de algunas herramientas algebraicas para obtener resultados dentro del capítulo 3, capítulo en el que además se revisan algunos objetos algebraicos asociados a los diagramas de un enlace. Finalmente, se concluye en el capítulo 4 con una exposición un poco informal de algunos resultados en física y su relación con la teoría de nudos, como el teorema de la estadística del espín y una aplicación de la ecuación de Yang Baxter.

Se recomienda para el lector una buena base de topología general y álgebra abstracta, la cual está cubierta (aunque no en su totalidad) en los apéndices de este trabajo. Cuando se presente algún concepto importante que requiera una revisión más profunda, se hará mención y se indicará una fuente recomendada.

1. NUDOS Y ENLACES

1.1. Teoría de nudos clásica

Antes de presentar de lleno las definiciones más robustas de la teoría de nudos, estableceremos algunas definiciones clásicas apelando un poco a la intuición. Ciertamente, esta teoría tiene sus bases establecidas en la topología, por lo que se invita al lector a consultar el Apéndice A (donde se hace una revisión de topología básica) si es su primer contacto con estas ramas.

1.1 Definición. Sea $K \subset \mathbb{R}^3$. Decimos que K es un *nudo* si existe un homeomorfismo $f : \mathbb{S}^1 \rightarrow K$, donde $\mathbb{S}^1$ es el círculo unitario, $\mathbb{S}^1 = \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 = 1\}$.

Un nudo, visto como curva en $\mathbb{R}^3$, es algo con lo que estamos familiarizados, pues son objetos que visualizamos constantemente, sin embargo podría parecer que estamos imponiendo una condición bastante fuerte al pedir que nuestra curva sea homeomorfa a $\mathbb{S}^1$, no obstante esta condición no limita tanto al nudo, pues simplemente con dar una función $\gamma : \mathbb{S}^1 \rightarrow \mathbb{R}^3$ continua e inyectiva, esta es un homeomorfismo en su imagen por el lema del mapeo cerrado, es decir, un encaje topológico. Por ello podemos concluir que la relación de homeomorfismo no es suficiente para establecer una igualdad entre nudos, esto motiva la siguiente definición:

1.2 Definición. Dos nudos K_1 y K_2 son equivalentes si existe un homeomorfismo $f : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ de forma que $f(K_1) = K_2$.

Aunque esta es la definición clásica y la que estudiaremos por ahora, en distintas fuentes existen otras nociones de equivalencia entre nudos.

Notemos que no estamos dando un homeomorfismo entre dos nudos, la condición de equivalencia es un homeomorfismo del ambiente en el que viven de tal forma que la imagen de uno de los nudos bajo esta aplicación es el otro. Realmente, esta condición va más orientada a la topología de un objeto de dimensión 3 que a la de un objeto de dimensión 1 (extenderemos esta discusión más adelante). Veamos ahora una forma estándar de representar nudos, usualmente lo haremos por medio

de *diagramas*, siendo el diagrama la proyección del nudo a $\mathbb{R}^2$. Dicha proyección se dibuja como una curva en la que, en cada intersección se borra parte de una curva dependiendo del nudo, para indicar que segmento pasa por encima o debajo como en la Figura 1.1. A cada intersección en la que ocurre este proceso se le denomina *cruce*.

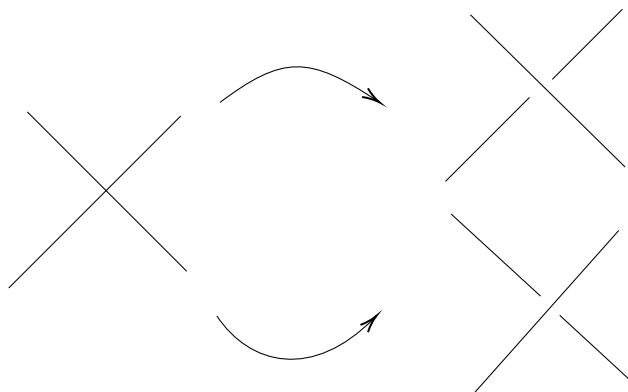


Figura 1.1. Cruces en diagramas para un nudo.

1.3 Ejemplo. El nudo más sencillo es conocido como nudo trivial o no nudo, siendo este $\mathbb{S}^1 \hookrightarrow \mathbb{R}^3$, es decir, $\mathbb{S}^1$ identificado como subconjunto de $\mathbb{R}^3$.

1.4 Ejemplo. Un nudo que puede mostrarse que no es equivalente al nudo trivial $\mathbb{S}^1$ es el nudo trébol 3_1 , el cual podemos ver en la Figura 1.2.

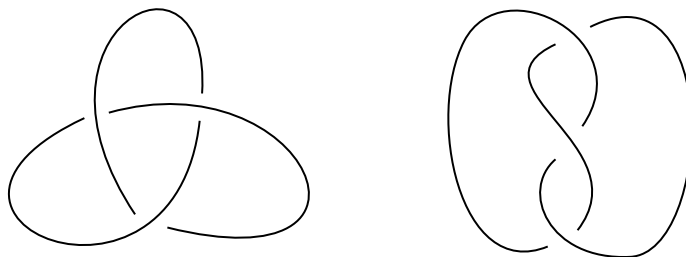


Figura 1.2. Dos representaciones del nudo trébol.

Para restringirnos un poco a los nudos que consideraremos de ahora en adelante, veamos las siguientes definiciones:

1.5 Definición. Un nudo se dice *poligonal* si es una curva cerrada simple que está formada por una cantidad finita de segmentos de recta.

Los diagramas de un nudo poligonal pueden dar problemas cuando consideramos sus cruces. Para eliminar algunas ambigüedades de estos diagramas veamos lo siguiente: si $\mathcal{P}(x, y, z) = (x, y, 0)$ es la proyección usada para construir el diagrama del nudo, entonces:

1.6 Definición. Dado un nudo K y $p \in \mathcal{P}(K)$ diremos que p es un *punto múltiple* si $\mathcal{P}^{-1}(p) \cap K$ tiene más de un elemento. A la cardinalidad de $\mathcal{P}^{-1}(p) \cap K$ la llamaremos *orden del punto p* .

1.7 Definición. Un nudo poligonal K está en una *posición regular* si

- i) Los únicos puntos múltiples de $\mathcal{P}(K)$ tienen orden dos.
- ii) Ningún punto doble de $\mathcal{P}(K)$ es preimagen de un vértice de K .

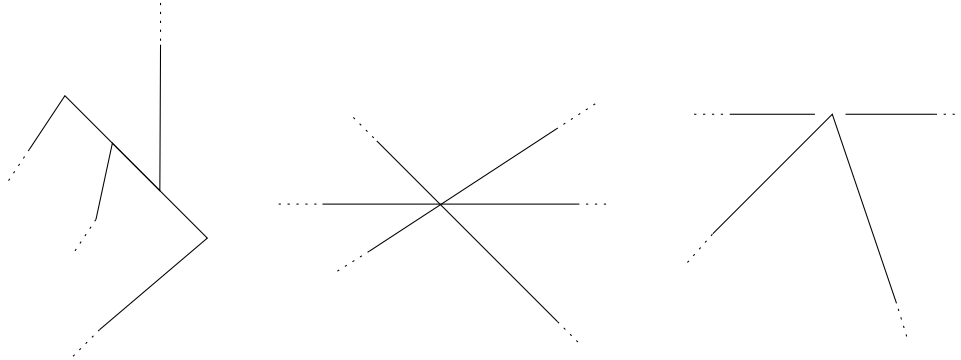


Figura 1.3. Cruces no permitidos de nudos en posición regular.

Como convención se asumirán que todos los nudos poligonales están en posición regular, pues todo nudo es equivalente a otro nudo en posición regular vía rotaciones.

1.8 Definición. Diremos que un nudo es *dócil* si es equivalente a un nudo poligonal, de lo contrario se dirá *salvaje*.

La mayoría de teoría sobre nudos se ha desarrollado para nudos dóciles, por lo que nos centraremos únicamente en estos de ahora en adelante.

1.9 Definición. El *número de palos*, $s(K)$, de un nudo K es el menor número de segmentos de recta necesarios para formar un nudo equivalente a K .

Con estas definiciones, probemos un pequeño resultado:

1.10 Proposición (Negami [27]). *Sea $c(K)$ el mínimo número de cruces para cualquier proyección de un nudo K , si K es no trivial, entonces*

$$\frac{5 + \sqrt{25 + 8(c(K) - 2)}}{2} \leq s(K).$$

Demostración. Consideremos la representación poligonal del nudo, y supongamos que esta viene dada por el mínimo número de segmentos necesarios para formar el nudo, es decir, $s(K)$ segmentos. Ahora lo rotaremos de forma que uno de estos segmentos sea paralelo al eje z , como en la Figura 1.4. Esto es posible hacerlo pues una rotación es un homeomorfismo de $\mathbb{R}^3$ a $\mathbb{R}^3$. Notemos que con esta modificación, en la proyección del nudo tendremos una arista menos.

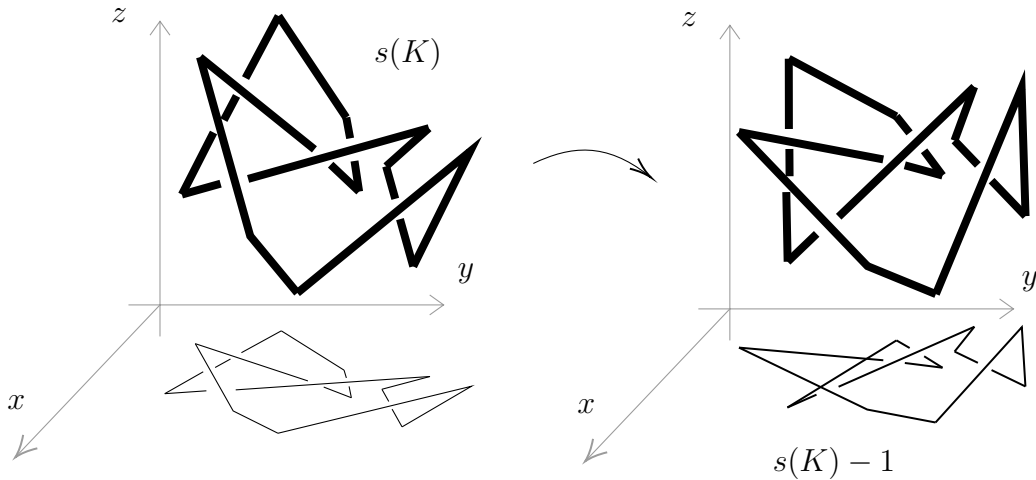


Figura 1.4. Rotación de un nudo en $\mathbb{R}^3$.

Por lo que, el diagrama del nudo rotado tendrá C cruces, tal que $c(K) \leq C$ y tendrá $s(K) - 1$ aristas (podría tener menos incluso, pero en ese caso la prueba es análoga e igual se puede llegar a la misma desigualdad). Como cada arista no puede cruzarse con sus aristas adyacentes a lo más tendremos $\frac{(s(K)-1)(s(K)-1-3)}{2}$ cruces, luego,

$$\begin{aligned} c(K) &\leq C \leq \frac{(s(K) - 1)(s(K) - 4)}{2} \\ 2c(K) &\leq (s(K))^2 - 5s(K) + 4 \\ 8c(K) &\leq 4(s(K))^2 - 20s(K) + 16 \\ 25 + 8(c(K) - 2) &\leq (2s(K) - 5)^2. \end{aligned}$$

Como para un nudo distinto del trivial ambos lados de la desigualdad son no negativos podemos tomar la raíz cuadrada y tendremos

$$\begin{aligned}\sqrt{25 + 8(c(K) - 2)} &\leq 2s(K) - 5 \\ \frac{5 + \sqrt{25 + 8(c(K) - 2)}}{2} &\leq s(K).\end{aligned}$$

■

Esto nos dice explícitamente que mientras más cruces tengamos, el nudo no puede ser equivalente a un polígono con pocos lados.

En general, podemos encontrar muchísimos nudos dentro de $\mathbb{R}^3$, de los cuales algunos están recopilados en [1, pp. 280-290]. Hay que resaltar que hasta las proposiciones más naturales pueden tener pruebas muy extensas si se desea dar una prueba completamente analítica de estas sin desarrollar muchas herramientas. Por ejemplo, dentro de los subconjuntos más naturales de $\mathbb{R}^2$ que podemos considerar, están los polígonos simples (aquellos sin autointersecciones) por lo que podríamos intentar a manera de ejercicio probar que estos son equivalentes al nudo trivial, una manera de hacerlo es la siguiente: consideremos la siguiente generalización del teorema de la curva de Jordan:

1.11 Teorema (Jordan-Schönflies). *Una curva cerrada simple C en un plano π separa π en dos regiones, y existe un homeomorfismo $\varphi : \pi \rightarrow \pi$ tal que $\varphi(C)$ es un círculo.*

Para una prueba elemental de este teorema puede consultarse [7]. Con la existencia de este homeomorfismo consideremos nuestro polígono P en $\mathbb{R}^2$ y al tomar su inclusión natural $P \subset \mathbb{R}^2 \hookrightarrow \mathbb{R}^3$ obtendremos la siguiente función

$$\tilde{\varphi}(x, y, z) = (\varphi(x, y), z)$$

la cual es un homeomorfismo $\tilde{\varphi} : \mathbb{R}^3 \rightarrow \mathbb{R}^3$, donde $\tilde{\varphi}(P)$ es un círculo. Esto muestra que nuestro polígono P visto como subconjunto de $\mathbb{R}^3$ es equivalente al nudo trivial.

Regresando a la discusión sobre el número de cruces de un nudo si, nos fijamos en los cruces podemos considerar un resultado bastante conocido:

1.12 Proposición. *Si el número de cruces en el diagrama de un nudo es 1 o 2, entonces el nudo es equivalente a un nudo trivial.*

Demostración. Para el caso en el que tenemos un solo cruce, podemos tomar una vecindad del punto en el que se da el cruce y podemos fijarnos en que, al ser el único cruce el resto del nudo será trozos de curvas de Jordan, de lo cual usando el Teorema 1.11 podemos reducir todos los posibles diagramas de los que se obtiene el cruce a los mostrados en la Figura 1.5. El diagrama de la izquierda en la figura no puede venir de un nudo, pues el diagrama indica que es una proyección de un conjunto disconexo. Por otro lado, para los otros dos es posible mostrar que estos son equivalentes al nudo trivial (veremos luego por qué).

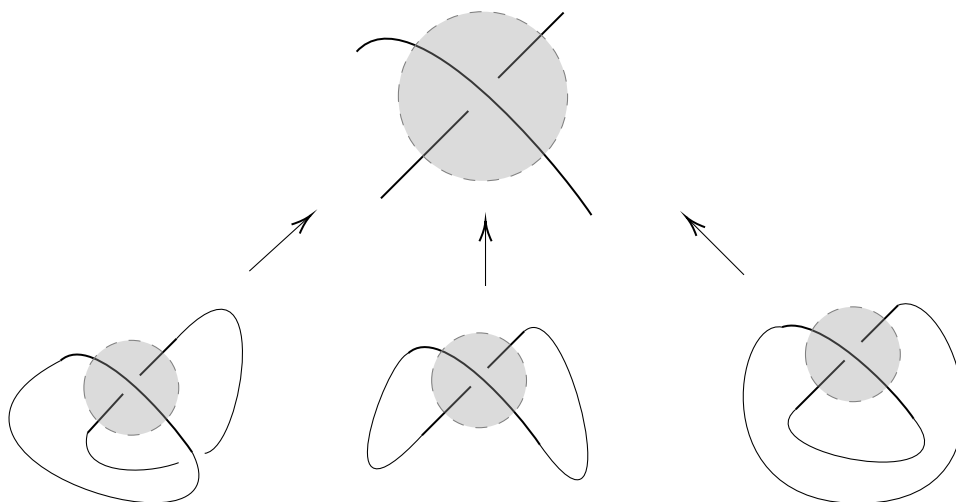


Figura 1.5. Diagramas de un nudo con un solo cruce

Finalmente puede hacerse una división en casos análoga para los diagramas con dos cruces y se puede concluir de la misma forma. ■

Ciertamente, para algunos podrían bastar estos argumentos visuales, pero es importante tener en cuenta que parte de la demostración se basa en una equivalencia entre el nudo trivial y un nudo con diagrama que parece tener forma de 8. Probar esa equivalencia de manera más rigurosa requiere de herramientas de las que no disponemos en este momento pero que vale la pena mencionar. Resulta que esta transformación se conoce como *Movimiento de Reidemeister I* y puede probarse que dos nudos equivalentes están relacionados por una secuencia finita de transformaciones similares, en concreto son tres transformaciones las que bastan para describir una equivalencia entre nudos. Se recomienda al lector llenar los detalles de esta discusión una vez se haya profundizado en la teoría detrás de los movimientos.

Como pudimos darnos cuenta, la mayoría de proposiciones que parecen inocentes en la teoría de nudos requieren de herramientas un poco técnicas para dar pruebas más analíticas, no obstante, siempre es bueno generar intuición con bocetos del problema. Para algunos de los ejemplos mencionados, se brindarán algunas de estas herramientas en secciones posteriores para rellenar algunos detalles en nuestras afirmaciones. Por ahora, en la misma línea de la teoría de nudos clásica veámos una forma de descomponer nudos, muy similar a como descomponemos números enteros.

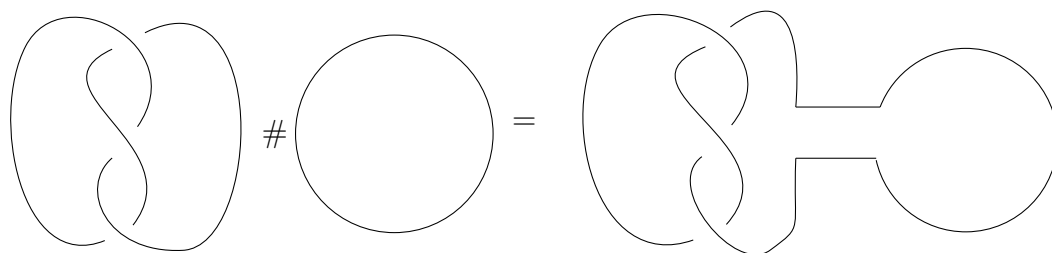


Figura 1.6. Idea de la composición de nudos.

Intuitivamente lo que deseáramos es tener dos nudos y definir alguna operación que involucre ambos para obtener otro nudo, tal y como la figura 1.6. Para definir bien esta operación, veamos antes algunas definiciones:

1.13 Definición. Sean X, Y espacios topológicos, A un subespacio cerrado de Y y $f : A \rightarrow X$ una aplicación continua. Además, si $\sim$ es la relación de equivalencia generada por los pares de la forma $(a, f(a))$, para todo $a \in A$. Se define el *espacio adjunto* del siguiente modo:

$$X \cup_f Y = X \sqcup Y / \sim .$$

Es decir, el espacio identificación de la unión disjunta de X e Y con la relación de equivalencia definida anteriormente.

Estos espacios son bastante generales y pueden derivarse muchas propiedades de ellos cuando X, Y son variedades de la misma dimensión, sin embargo no indagaremos tanto en sus propiedades. Para una discusión más detallada de estos espacios puede consultarse [5].

1.14 Definición. Sea M una n -variedad. Diremos que una bola de coordenadas $B \subseteq M$ es una *bola regular de coordenadas*, si existe una vecindad B' de $\overline{B}$ y un homeomorfismo $\varphi : B' \rightarrow B_{r'}(x) \subseteq \mathbb{R}^n$ tal que mapea B a $B_r(x)$ y $\overline{B}$ a $\overline{B}_r(x)$ para algún $r' > r > 0$ y $x \in \mathbb{R}^n$.

1.15 Definición. Sean M_1 y M_2 n -variedades conexas con $B_i \subset M_i$ bolas regulares de coordenadas. Los subespacios $M'_i = M_i \setminus B_i$ serán n -variedades con frontera, tales que las fronteras son homeomorfas a $\mathbb{S}^{n-1}$. Si $f : \partial M'_2 \rightarrow \partial M'_1$ es cualquier homeomorfismo definimos:

$$M_1 \# M_2 = M'_1 \cup_f M'_2$$

como la *suma conexa* de M_1 y M_2 .

Ahora, hay un detalle omitido: nótese que la definición de $M_1 \# M_2$ depende de un homeomorfismo f y la elección de bolas regulares de coordenadas. Un resultado altamente no trivial es, en efecto, mostrar que la suma conexa de dos variedades nos devuelve a lo más dos variedades salvo homeomorfismos,¹ las dos posibilidades corresponden a los casos en los que el homeomorfismo f preserva o no la orientación de la frontera. Por ello, en este trabajo no se presentarán variedades en las que esta definición sea ambigua o directamente se indicará a cual de las dos opciones se está refiriendo la suma conexa. En concreto, el que solo existan dos posibles sumas conexas se soluciona dotando de orientación a las variedades en que aplicaremos esta operación, particular para nudos, solo se consideran las sumas conexas en las que la orientación de ambos nudos coincide. A continuación, consideremos la siguiente como una definición de orientación para nudos.

1.16 Definición. Una *orientación* es una elección para viajar a través del diagrama del nudo. Esta dirección es denotada en el diagrama por flechas. A un nudo provisto de una orientación le llamaremos *nudo orientado*, vistas en la Figura 1.7.

Con este detalle aclarado, es posible mostrar que si K_1 es un nudo arbitrario y K_2 es el nudo trivial, entonces $K_1 \# K_2 = K_1$, donde en este caso el símbolo igual indica equivalencia de nudos, vistas en la Figura 1.8.

Esto nos invita a pensar en los nudos como una estructura algebraica y las siguientes definiciones:

¹Este resultado depende de otra proposición conocida como el *Annulus theorem*, del cual sus primeras pruebas datan de poco más de 100 años y sus demostraciones son distintas para dimensiones bajas, es decir para variedades de dimensión 1 a 4.

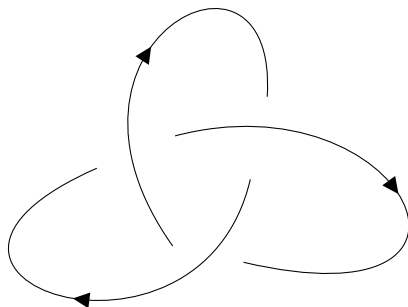


Figura 1.7. Nudo trébol orientado

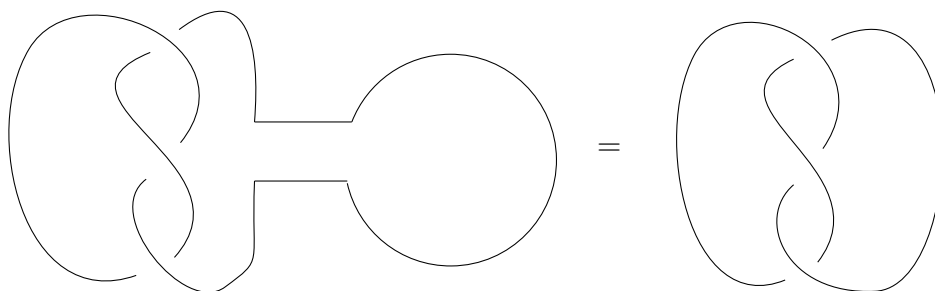


Figura 1.8. Operación con el nudo trivial.

1.17 Definición. Un nudo K es llamado *compuesto* si puede ser expresado como la suma conexa de dos nudos, donde ninguno de los nudos por los cuales está formado son el nudo trivial, esto es $K = K_1 \# K_2$ donde K_1, K_2 son no triviales.

A los nudos K_1 y K_2 , se les conoce como nudos factores. Finalmente se definen los análogos a los números primos en nudos.

1.18 Definición. Un nudo que no puede escribirse como suma conexa de cualesquiera dos nudos no triviales es llamado *nudo primo*.

Aunque no ocurre en general, usualmente la suma conexa de nudos estará bien definida, es decir, los dos posibles nudos que resultan serán equivalentes. Los nudos en los que ocurre esto se conocen como nudos *invertibles* y uno de ellos es el trébol 3_1 el cual es un nudo primo también y un ejemplo de un nudo no invertible es uno de

los nudos con 8 cruces, más concretamente el 8_{17} . En la Figura 1.9 pueden observarse dos posibles sumas conexas de el nudo 8_{17} consigo mismo.

A día de hoy se han clasificado completamente los nudos primos de hasta 19 cruces, esto haciendo uso de herramientas computacionales. Para más detalles sobre las técnicas usadas en la clasificación de nudos hasta 19 cruces puede consultarse [6].

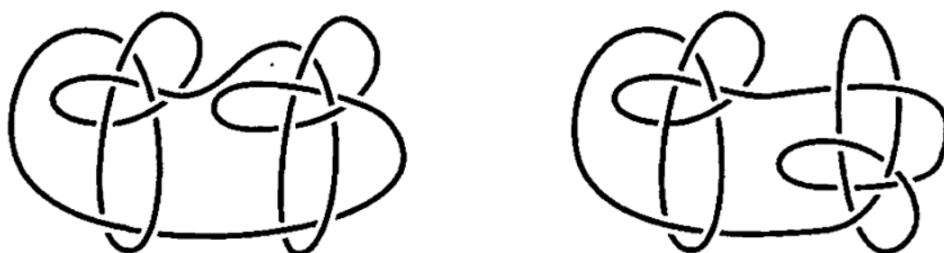


Figura 1.9. Dos nudos distintos con mismos factores. Fuente: Tomada de [1, p. 12]

1.19 Definición. Un *enlace*, es una unión disjunta de una cantidad finita de nudos. Es decir si $K_1, K_2, \dots, K_n$ nudos entonces $L = \sqcup_{i=1}^n K_i$ es un enlace y a cada uno de los nudos K_i se les conoce como *componentes del enlace*.

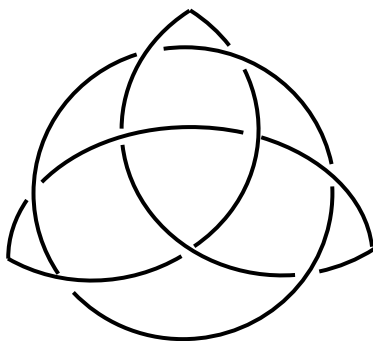


Figura 1.10. Triqueta con un círculo

1.20 Ejemplo. El enlace de la figura 1.10 es la triqueta con un círculo puede ser vista como la unión disjunta del nudo trébol y el nudo trivial.

1.21 Definición. Un enlace en el que todas sus componentes son nudos poligonales se conoce como *enlace poligonal*.

Se define así, una noción de equivalencia entre enlaces poligonales:

1.22 Definición. Sea u un segmento de recta perteneciente a uno de los nudos que componen un enlace poligonal $L \subset \mathbb{R}^3$. Sea Δ^2 un triángulo en $\mathbb{R}^3$ cuya frontera consiste en tres segmentos de recta denotados por u, v y w de modo que $\Delta^2 \cap L = u$. La curva poligonal L' , definida como $L' = (L - u) \cup v \cup w$, es un nuevo enlace poligonal en $\mathbb{R}^3$. Así decimos que L' ha sido obtenido desde L vía un Δ -movimiento. De forma opuesta, decimos que L se obtiene a partir de L' por un Δ^{-1} -movimiento. Vistas en la figura 1.11.

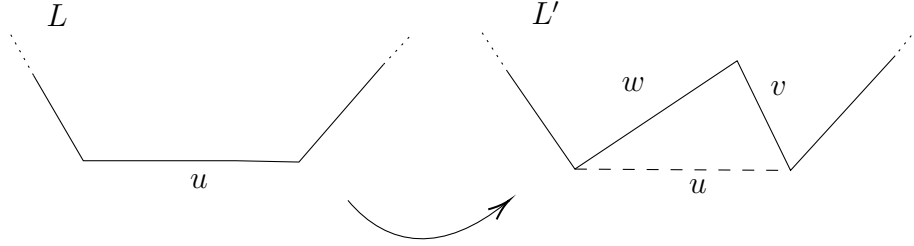


Figura 1.11. Δ -Movimiento.

1.23 Definición. Dos enlaces poligonales se dicen ser Δ -equivalentes si uno puede ser obtenido del otro a través de una sucesión finita de Δ -movimientos y Δ^{-1} -movimientos.

Con esto tenemos las herramientas necesarias para enunciar uno de los teoremas principales de la teoría de nudos. El siguiente teorema data de 1927 y es de suma importancia para la teoría de nudos, pues simplifica el análisis de equivalencias.

1.24 Teorema (Reidemeister [29]). *Dos diagramas tienen enlaces Δ -equivalentes si y sólo si los diagramas están relacionados a través de una sucesión finita de movimientos de Reidemeister $\mathcal{R}_i$ para $i = 1, 2, 3$ y deformación del plano del diagrama. Vistas en la Figura 1.12.*

La relación de que dos enlaces sean Δ -equivalentes está ampliamente relacionada con la relación de que dos nudos sean equivalentes. Basta con dar una sucesión de estos movimientos aplicados al diagrama del nudo que lo transformen en el diagrama de otro nudo para mostrar que ambos son equivalentes. Si dicha sucesión existe diremos que existe una *isotopía de ambiente* entre ambos nudos o que los nudos son *isotópicos*. Si la sucesión de movimientos solo contiene movimientos del tipo $\mathcal{R}_2$ y $\mathcal{R}_3$, a la equivalencia la llamaremos *isotopía regular*.

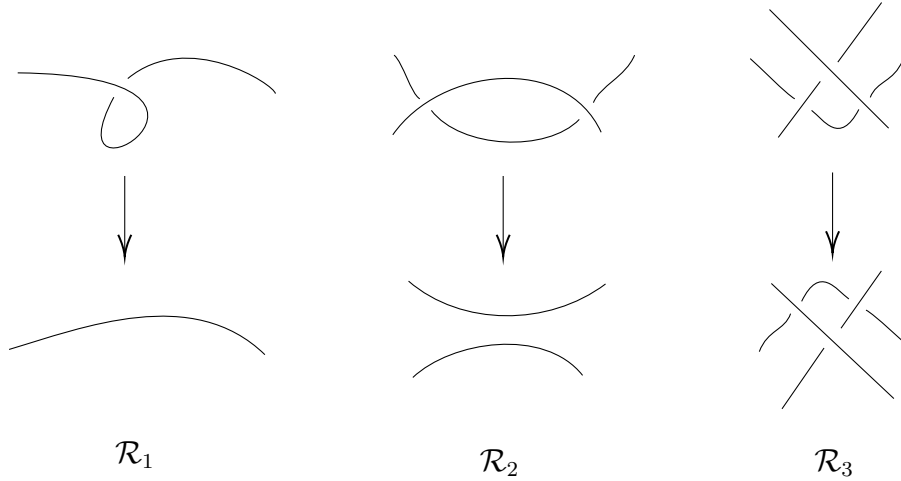


Figura 1.12. Movimientos de Reidemeister

1.2. Teoría de nudos en la 3-esfera

Ahora presentaremos un enfoque más moderno de la teoría de nudos, pues en la definición de nudo podemos reemplazar $\mathbb{R}^3$ por una 3-variedad y obtener resultados análogos. Esto se debe a que al mantener la codimensión tenemos suficiente libertad para tener más de un nudo, pero no tanta libertad para que nuestra teoría se trivialice. Para indagar un poco en la teoría de 3-variedades, primero se presentan algunas definiciones generales.

1.25 Definición. El conjunto $\{v_0, v_1, \dots, v_k\} \subset \mathbb{R}^n$ se dice *afínmente independiente* o *están en posición general* si el conjunto $\{v_1 - v_0, v_2 - v_0, \dots, v_k - v_0\}$ es linealmente independiente.

1.26 Definición. Dado un conjunto $\{v_0, v_1, \dots, v_k\}$ afínmente independiente de $k + 1$ puntos, el *simplejo* generado por ellos se denota por $[v_0, \dots, v_k]$ y se define como:

$$[v_0, \dots, v_k] = \text{conv}\{v_0, \dots, v_k\} = \left\{ \sum_{i=0}^k t_i v_i \mid \sum_{i=0}^k t_i = 1, \text{ con } t_i \geq 0 \right\}.$$

Cada uno de los puntos v_i se denomina *vértice del simplejo*. El entero k (uno menos que el número de vértices) se denomina *dimensión*; un simplejo de dimensión k se suele denominar *k-simplejo*.

1.27 Definición. Si $\sigma = [v_0, \dots, v_k]$ es un k -simplejo, el simplejo $\text{conv}\{w_0, \dots, w_\ell\}$ es llamado *cara* de σ si $\{w_0, \dots, w_\ell\} \subset \{v_0, \dots, v_k\}$.

Notemos que en esta definición la dimensión de la cara será $\ell \leq k$. En el caso $\ell < k$ escribiremos $\tau < \sigma$ y en particular cuando $\ell = 0$ tendremos los vértices y cuando $\ell = 1$ a la cara la llamaremos *arista*.

1.28 Ejemplo. En $\mathbb{R}^3$ podemos construir 4 simplejos los cuales están ilustrados en la figura 1.13: el 0-simplejo que es un punto, el 1-simplejo que es un segmento de recta, el 2-simplejo conocido como triángulo y el 3-simplejo, más conocido como tetraedro.

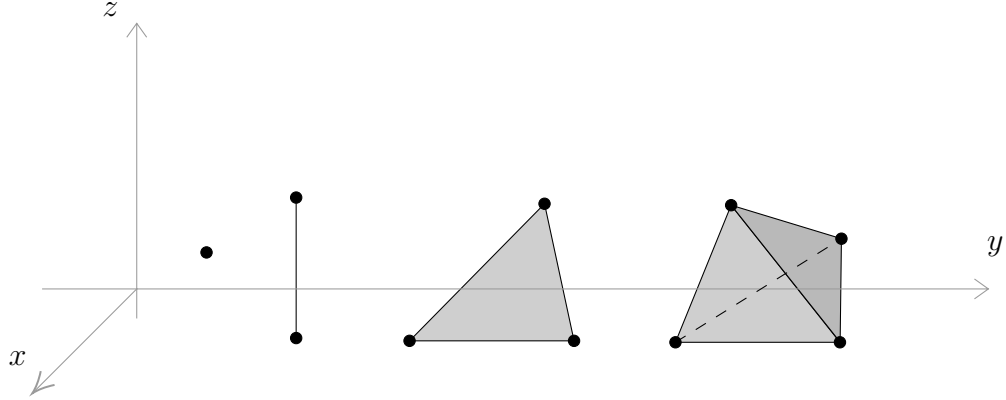


Figura 1.13. Simplejos en $\mathbb{R}^3$.

1.29 Ejemplo. El n -simplejo estandar es un subconjunto del espacio $\mathbb{R}^n$ definido por $\Delta^n = \{(t_1, \dots, t_n) \in \mathbb{R}^n \mid \sum_{i=1}^n t_i \leq 1, t_i \geq 0 \forall i\}$. En este caso los vértices son la base canónica $\{e_i\}_{1 \leq i \leq n}$ y el origen de $\mathbb{R}^n$.

1.30 Definición. Una colección K de simplejos en $\mathbb{R}^n$ es llamada *complejo simplicial* si cumple las siguientes condiciones:

1. Si $\sigma \in K$ y $\tau < \sigma$ entonces $\tau \in K$.
2. Si $\sigma, \tau \in K$ entonces $\sigma \cap \tau < \sigma$ y $\sigma \cap \tau < \tau$.
3. K es localmente finita, esto es, dado $x \in \sigma \in K$, existe una vecindad de x en $\mathbb{R}^n$ que intersecta un número finito de elementos de K .

1.31 Definición. Dos complejos simpliciales K, L se dicen *isomorfos* si existe una biyección que preserve caras entre ellos.

1.32 Definición. Dado un complejo simplicial K se define el *poliedro de K* como:

$$|K| = \bigcup_{\sigma \in K} \sigma.$$

1.33 Definición. Un complejo K' es una *subdivisión* de K dado que $|K'| = |K|$ y cada simplejo de τ en K' es subconjunto de algún simplejo de $\sigma \in K$. En este caso escribimos $K' \prec K$.

1.34 Definición. Dado un complejo simplicial K y un simplejo $\sigma \in K$, la *estrella* de σ en K es el subcomplejo simplicial denotado por $\text{St}(\sigma, K)$ definido como:

$$\text{St}(\sigma, K) = \{\tau \in K \mid \exists \eta \text{ tal que } \sigma, \tau < \eta\},$$

donde además $\text{st}(\sigma, K) = |\text{St}(\sigma, K)|$.

1.35 Definición. Dado un complejo simplicial K y un simplejo $\sigma \in K$, el *enlace* de σ en K es el subcomplejo simplicial denotado por: $\text{Lk}(\sigma, K)$ definido como:

$$\text{Lk}(\sigma, K) = \{\tau \in \text{St}(\sigma, K) \mid \tau \cap \sigma = \emptyset\},$$

donde además $\text{lk}(\sigma, K) = |\text{Lk}(\sigma, K)|$.

Ahora nos fijaremos en algunas funciones entre complejos simpliciales. Recordemos que una función lineal en un espacio vectorial es una función tal que $f(\alpha v + \beta w) = \alpha f(v) + \beta f(w)$, para todos α, β en el campo en el que está definido el espacio.

1.36 Definición. Dados K, L complejos simpliciales una *aplicación simplicial* es una función $f : |K| \rightarrow |L|$ tal que para todo $\sigma \in K$, la restricción $f|_\sigma$ es lineal y $f|_\sigma(\sigma)$ es un simplejo en L .

Algunas veces se abusa de la notación convenientemente y nos referiremos a las aplicaciones simpliciales como funciones $f : K \rightarrow L$.

1.37 Definición. Dos poliedros son *PL-homeomorfos* y se denota por $K \cong_{PL} L$ si existen subdivisiones $K' \prec K$ y $L' \prec L$ tales que K' es isomorfo a L' .

1.38 Definición (*Variedad combinatoria*). Una *n-variedad combinatoria* es un complejo simplicial tal que el enlace de cada p -simplejo es PL-homeomorfo a la frontera de un $(n - p)$ -simplejo o a un $(n - p - 1)$ -simplejo.

1.39 Definición. Una *triangulación* de un espacio topológico X consiste un complejo simplicial K y un homeomorfismo $t : |K| \rightarrow X$.

1.40 Definición (*Variedad PL*). Una *n-variedad PL* es una *n-variedad topológica* M con una triangulación $t : |K| \rightarrow X$, donde K es una *n-variedad combinatoria*.

Autores como Lickorish [20] definen un enlace como un subconjunto de $\mathbb{S}^3$ compuesto por m curvas cerradas simples, lineales a trozos y un nudo como el caso cuando $m = 1$. Donde por curva lineal a trozos dentro de $\mathbb{S}^3$ nos referimos a una 1-variedad PL. La teoría de variedades PL permite restringir a nuestras curvas lo suficiente para no obtener nudos salvajes ya que la condición (3) en la definición 1.30 descarta los nudos se ven como el descrito en la figura 1.14 ya que hay un punto que, claramente, dará problemas.

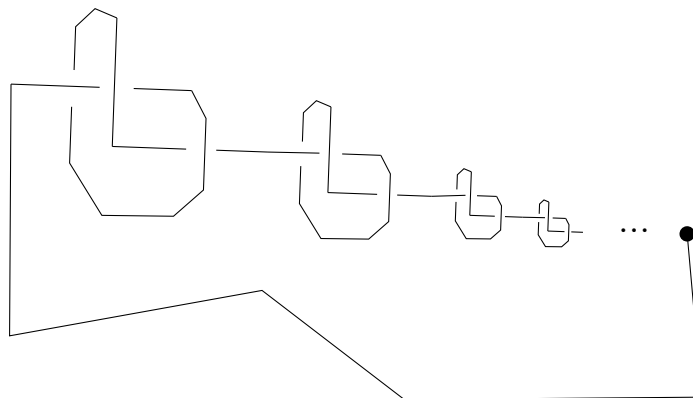


Figura 1.14. Nudo salvaje

La condición de ser variedad PL puede sustituirse pidiendo que nuestras curvas sean suaves (es decir 1-variedades suaves encajadas en $\mathbb{S}^3$). Si se desea indagar más sobre la teoría de variedades PL, puede consultarse [10] y [14].

Los dos ambientes más populares para encajar nuestros nudos son $\mathbb{S}^3$ y $\mathbb{R}^3$. Sin embargo no son las únicas variedades de dimensión 3 en las que se hace teoría de nudos; en [11] han estudiado nudos en el espacio proyectivo $\mathbb{R}P^3$. No obstante existen algunas ventajas de trabajar nudos en $\mathbb{S}^3$ por encima de $\mathbb{R}^3$, una de ellas es que la compacidad de la 3-esfera nos da más información de la variedad y que, realmente, es un espacio muy parecido a $\mathbb{R}^3$, algo que estamos acostumbrados a ver, ya la 3-esfera es la compactificación del espacio euclideo tridimensional. Otra ventaja con la que no entraremos en muchos detalles es que existe más información de algunos invariantes en $\mathbb{S}^3$, como los grupos de homotopía superior.

Para seguir indagando en el enfoque de la teoría de nudos en la 3-esfera necesitaremos una definición auxiliar:

1.41 Definición. Si X, Y son espacios topológicos. Dos homeomorfismos $f, g : X \rightarrow Y$ se dicen *isotópicos* si existe una aplicación continua

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

tal que $H(x, 0) = f(x)$ y $H(x, 1) = g(x)$ y $H_t(x) = H(x, t)$ es un homeomorfismo para cada $t \in [0, 1]$ fijo. A la función H se le conoce como *isotopía*.

Algunos autores dan diferentes nociones de equivalencia entre nudos, dos de las más estandar son las siguientes:

- i) Dos nudos K_1, K_2 en $\mathbb{S}^3$ son equivalentes si existe un homeomorfismo que preserva orientación $h : \mathbb{S}^3 \rightarrow \mathbb{S}^3$ tal que $h(K_1) = K_2$.
- ii) Dos nudos K_1, K_2 en $\mathbb{S}^3$ son equivalentes si existe una isotopía entre $\text{Id}_{\mathbb{S}^3}$ y $h : \mathbb{S}^3 \rightarrow \mathbb{S}^3$ un homeomorfismo que preserva orientación tal que $h(K_1) = K_2$. En este caso se dice que K_1 y K_2 tienen el mismo *tipo de isotopía*.

Sin embargo estas dos definiciones de equivalencia de nudos, son equivalentes y la prueba de ello viene de un resultado bastante general que puede ser encontrado en [24, p. 2], en particular el caso $n = 3$ del teorema del artículo de Livingston establece lo siguiente:

1.42 Teorema. *Todo homeomorfismo que preserva orientación $h : \mathbb{S}^3 \rightarrow \mathbb{S}^3$ es isotópico a la identidad en $\mathbb{S}^3$.*

1.3. Teoría general de nudos

Al principio de este capítulo se hizo mención de que, la teoría de nudos realmente pone más atención a la topología de una variedad de dimensión 3 que a la topología de los objetos de dimensión 1. Posteriormente se hace mención de que esta relación entre las dimensiones de las variedades es importante. Esto se debe a lo siguiente: algunos autores como Rolfsen [31] definen un nudo como la imagen K de un encaje topológico $f : \mathbb{S}^k \rightarrow \mathbb{S}^n$ y en concreto si se mantiene la definición de equivalencia de nudos, obtendremos que por ejemplo, para el caso $(n, k) = (2, 1)$ el teorema de Jordan-Schönflies nos asegura que solamente hay un nudo. Con un poco más de trabajo y robusteciendo la definición de equivalencia puede mostrarse un análogo para $(n, k) = (4, 1)$.

Así que un buen punto para intentar generalizar nuestra teoría de nudos será mantener la codimensión, en concreto, considerando el par $(n, k) = (4, 2)$. Y tal como se hizo en la teoría de nudos clásica antes de considerar los nudos como encajes en esferas pensemos a los nudos como superficies (variedades de dimensión 2) en $\mathbb{R}^4$. Resulta que la teoría para este caso está desarrollada, al considerar una isotopía de ambiente en superficies como nuestra relación de equivalencia. Existen análogos de los movimientos de Reidemeister llamados *movie moves* de los cuales podemos ver algunos en la figura 1.15. Teniendo ahora 15 movimientos tales que si existe una sucesión de ellos, tendremos una isotopía ambiente entre superficies.

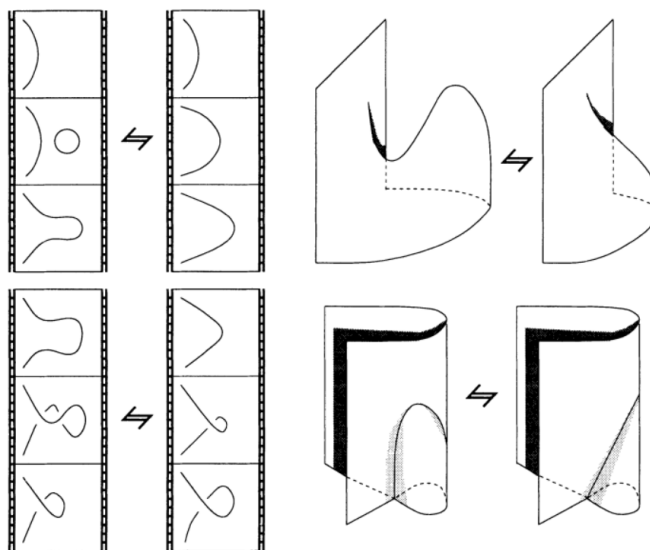


Figura 1.15. Movie moves. Fuente: Tomada de [9, p. 46].

Donde por isotopía ambiente entre superficies nos referimos a lo siguiente:

1.43 Definición. Dos superficies $f_i : F \rightarrow \mathbb{R}^4$, $i = 0, 1$ son llamadas *ambientalmente isotópicas* si existe una isotopía $H : \mathbb{R}^4 \times I \rightarrow \mathbb{R}^4$ tal que $H(x, 0) = x$ para todo $x \in \mathbb{R}^4$ y $H(f_0(a), 1) = f_1(a)$ para todo $a \in F$.

En [8] se da la completa clasificación de estos movimientos y finalmente en [9] se considera una definición bastante general de nudo, la cual encaja con nuestra discusión.

1.44 Definición. Un *nudo generalizado* es un encaje topológico de una variedad cerrada de codimensión 2, $K : M^{n-2} \rightarrow N^n$.

Donde usualmente $N^n = \mathbb{R}^n$ y $M^{n-2} = \mathbb{S}^{n-2}$, con la notación de la definición.

2. TEOREMA DE SEIFERT-VAN KAMPEN

En este capítulo se desarrollan las herramientas necesarias para la construcción de uno de los teoremas más importantes dentro de la topología algebraica. El teorema de Seifert-Van Kampen es una de las técnicas más potentes para el cálculo de grupos fundamentales, pues nos permite calcular el grupo fundamental de un espacio topológico en función de subespacios de este, con algunas restricciones, siendo también el grupo fundamental un functor, en el que centraremos nuestra atención desde ahora, para definir posteriormente el grupo de un nudo, introduciendo así, nuestro primer invariante algebraico. El desarrollo de estas herramientas tiene como objetivo dar lugar en el siguiente capítulo a la primera prueba explícita dentro de este trabajo de la existencia de nudos no triviales. Hay que destacar que muchas de las pruebas en este capítulo están basadas en [19] y [26], pues son textos referentes en la materia.

2.1. Equivalencia Homotópica

Como punto de partida, introduciremos una relación de equivalencia entre aplicaciones continuas de un espacio topológico a otro, dando posteriormente un vistazo a la noción de equivalencia homotópica entre dos espacios topológicos. Precisamente estos conceptos son fundamentales en la clasificación de espacios topológicos pues veremos que en efecto, son más generales que la relación de homeomorfismo. A lo largo de todo el capítulo tomaremos como convención $I = [0, 1]$.

2.1.1. Homotopía

2.1 Definición (Homotopía Libre). Dados X, Y dos espacios topológicos, dos aplicaciones continuas $f, g : X \rightarrow Y$ se dicen *homotópicas* o *libremente homotópicas* si existe una función continua $H : X \times I \rightarrow Y$ tal que:

$$H(x, 0) = f(x), \quad H(x, 1) = g(x) \quad \forall x \in X.$$

A la aplicación H se le conoce como *homotopía libre* o simplemente *homotopía* entre f y g . En este caso, escribimos $f \simeq g$.

El término “homotopía libre” hace referencia a que estamos hablando de una deformación en donde los objetos que estamos involucrando tienen total libertad en los valores que toman. Es muy común que las aplicaciones tengan ciertas restricciones y por ello conviene hacer la siguiente distinción:

2.2 Definición. Una homotopía $H : X \times I \rightarrow Y$ entre f y g se dice *estacionaria en A* o *relativa a A* si

$$H(x, t) = f(x) \quad \forall x \in A, \forall t \in I.$$

Esto es, para cada t la función $H_t(x) = H(x, t)$ coincide con f en A . En este caso diremos que f y g son *homotópicas relativas al conjunto A* y lo denotaremos por $f \simeq_A g$.

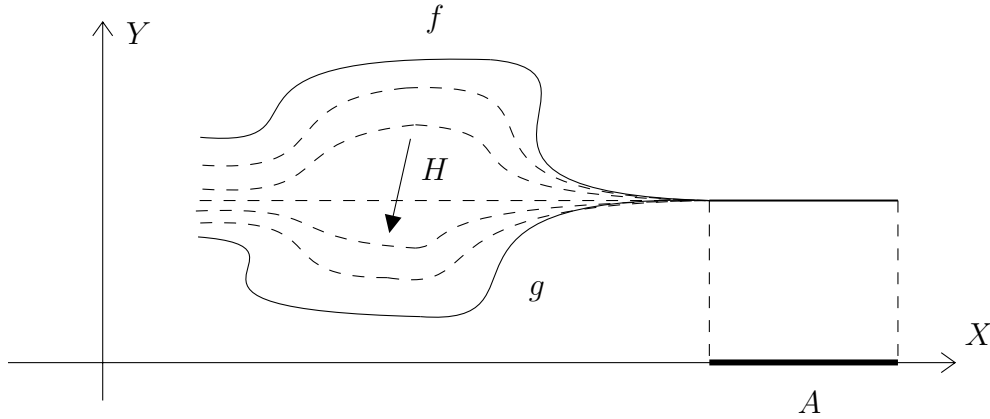


Figura 2.1. Homotopía relativa a un conjunto A .

Notemos que en este caso es necesario que $g|_A = H_1|_A = f|_A$, así que podemos pensar que dos funciones para ser homotópicas relativamente en A , deben coincidir al menos en A .

Geométricamente, una homotopía se interpreta como en la Figura 2.1. Además las relaciones $\simeq$ y $\simeq_A$ son de equivalencia, para esto presentaremos una prueba.

2.3 Proposición. Dados X, Y espacios topológicos y $C(X, Y)$ el conjunto de todas las aplicaciones continuas de X a Y , las relaciones $\simeq$ y $\simeq_A$ en $C(X, Y) \times C(X, Y)$ son de equivalencia.

Demostración. Haremos la prueba para la relación $\simeq$, pues $\simeq_A$ es un caso particular. Esta relación es reflexiva pues podemos considerar $H(x, t) = f(x) \forall x$ y al ser f continua tendremos la continuidad de H , obteniendo una homotopía entre f y f . Si $f \simeq g$, existe $H : X \times I \rightarrow Y$ homotopía entre f y g , y definiendo $\overline{H}(x, t) = H(x, 1 - t)$ se obtiene una función continua que será nuestra homotopía entre g y f , por ello la relación es simétrica. Finalmente para probar que es transitiva supongamos $f \simeq g$ y $g \simeq h$, donde H una homotopía entre f y g y K una homotopía entre g y h , con esto definimos:

$$F(x, t) = \begin{cases} H(x, 2t) & \text{si } t \in [0, \frac{1}{2}] \\ K(x, 2t - 1) & \text{si } t \in [\frac{1}{2}, 1] \end{cases}$$

Ahora, esta función está bien definida pues se tiene que $H(x, 2t)$ y $K(x, 2t - 1)$ coinciden en el cerrado $X \times \{\frac{1}{2}\}$ y por el Lema del pegado (Proposición A.18) se sigue que F es continua, teniendo así que F es una homotopía entre f y h . ■

Algo que podemos notar es que cuando es posible dotar a $C(X, Y)$ de una topología (usualmente se toma la topología compacto-abierta) cada homotopía puede ser vista como una camino entre dos funciones, más precisamente, si H es una homotopía entre f y g se tiene que $\gamma : I \rightarrow C(X, Y)$ es un camino en $C(X, Y)$ definido por $\gamma(t) = H_t$ y las clases de equivalencia generadas por la relación de equivalencia $\simeq$ serán componentes conexas por caminos de $C(X, Y)$. Para más detalles sobre la topología compacto-abierta puede revisarse [4, Cap. 7] en donde además se exponen algunas aplicaciones de este enfoque.

2.1.2. Tipo de homotopía

Ahora presentaremos una equivalencia topológica más débil que la condición de homeomorfismo, para enriquecer el estudio de las homotopías.

2.4 Definición. Una aplicación continua $f : X \rightarrow Y$ es llamada *equivalencia homotópica* cuando existe $g : Y \rightarrow X$ continua tal que $g \circ f \simeq \text{Id}_X$ y $f \circ g \simeq \text{Id}_Y$. En este caso, decimos que g es el *inverso homotópico* de f .

2.5 Definición. Si existe una equivalencia homotópica entre dos espacios topológicos X, Y se dice que X tiene el *mismo tipo de homotopía* que Y y escribiremos $X \equiv Y$.

Una consecuencia de esta definición es que si tomamos dos espacios homeomorfos estos tendrán el mismo tipo de homotopía pues directamente tendremos funciones continuas f, g entre dos espacios tales que su composición es igual a la identidad haciendo ambas composiciones, por lo que basta tomar una homotopía constante en una variable. Escrito de manera más precisa tendremos:

2.6 Proposición. *Sean X, Y dos espacios topológicos homeomorfos, entonces estos tienen el mismo tipo de homotopía.*

Veamos algunos ejemplos no tan estándar de espacios con el mismo tipo de homotopía.

2.7 Ejemplo. La banda de Möbius tiene el mismo tipo de homotopía que un cilindro. La idea para ver esto es tratar de retraer la banda a una copia de $\mathbb{S}^1$ dentro de ella misma, tal y como podemos ver en la Figura 2.2. Para probar esto vamos considerar a la banda de Möbius como un espacio cociente $M = I^2 / \sim$ donde $\sim$ es la relación de equivalencia generada por $(0, y) \sim (1, 1 - y)$, para todo $y \in I$. De esta manera, si $q : I^2 \rightarrow M$ es la aplicación cociente, definimos la función f como

$$f : M \rightarrow \mathbb{S}^1,$$

$$[(x, y)] \mapsto f[(x, y)] = e^{2\pi i x},$$

veremos que esta es continua pues $(f \circ q)(x, y) = e^{2\pi i x}$ es continua.

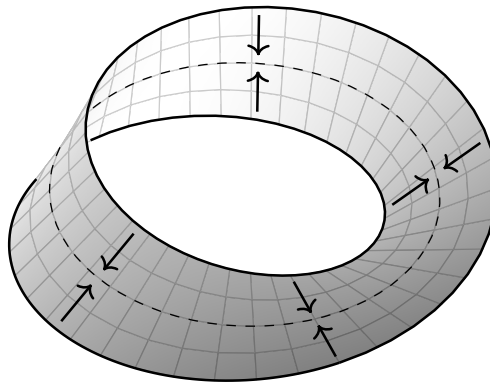


Figura 2.2. Equivalencia de la banda de Möbius

La función $g_o(t) = \left[\left(t, \frac{1}{2} \right) \right]$ es continua y pasando para el cociente, Proposición A.21

obtendremos una función continua $g : \mathbb{S}^1 \rightarrow M$ tal que $g(e^{2\pi it}) = [(t, \frac{1}{2})]$. En otras palabras, para $\varepsilon(t) = e^{2\pi it}$ se tiene

$$\begin{array}{ccc} I & \xrightarrow{g \circ \varepsilon} & M \\ \varepsilon \downarrow & \nearrow g & \\ \mathbb{S}^1 & & \end{array}$$

luego, para $z \in \mathbb{S}^1$ arbitrario se expresa $z = e^{2\pi it}$ obteniendo

$$(f \circ g)(z) = f(g(z)) = f(g(e^{2\pi it})) = f[(t, 1/2)] = e^{2\pi it} = z,$$

y así $f \circ g = \text{Id}_{\mathbb{S}^1}$. Por otro lado $(g \circ f)[(x, y)] = [(x, \frac{1}{2})]$, pues para ver que efectivamente $g \circ f \simeq \text{Id}_M$, se considera función

$$H([(x, y)], t) = \left[\left(x, \frac{t}{2} + (1-t)y \right) \right],$$

siendo esta una homotopía entre $g \circ f$ y Id_M , pues la continuidad se deriva de que la función

$$\begin{aligned} Q : I^2 \times I &\rightarrow M \times I \\ ((x, y), t) &\mapsto ([(x, y)], t) \end{aligned}$$

es un mapeo cociente y $H \circ Q$ es continua, por lo que H es continua. Finalmente nuestro trabajo se reduce a probar que $\mathbb{S}^1$ tiene el mismo tipo de homotopía que un cilindro, digamos $\mathbb{S}^1 \times I$. Para esto se consideran las funciones

$$\begin{aligned} h : \mathbb{S}^1 \times I &\rightarrow \mathbb{S}^1 & \iota : \mathbb{S}^1 &\hookrightarrow \mathbb{S}^1 \times I \\ (z, t) &\mapsto z & z &\mapsto (z, 0) \end{aligned}$$

y en efecto estas verifican ser inversos homotópicos, por lo que $M \equiv \mathbb{S}^1 \equiv \mathbb{S}^1 \times I$.

Este ejemplo previo muestra que la equivalencia homotópica es más débil que la relación de homeomorfismo ya que la banda de Möbius no es homeomorfa al cilindro y esto se puede ver al considerar la curva cerrada ∂M marcada en rojo en la figura 2.3. Notemos que un homeomorfismo hipotético φ mapearía ∂M a otra curva cerrada sin autointersecciones dentro del cilindro como las mostradas a la derecha en la figura 2.3. Ahora esto es un problema, pues toda curva cerrada simple separa al cilindro en dos componentes conexas y la banda de Möbius no tiene esa

propiedad, en otras palabras $M \setminus \partial M$ es conexo pero $\varphi(M \setminus \partial M)$ no lo es, de ahí que estos espacios no pueden ser homeomorfos.

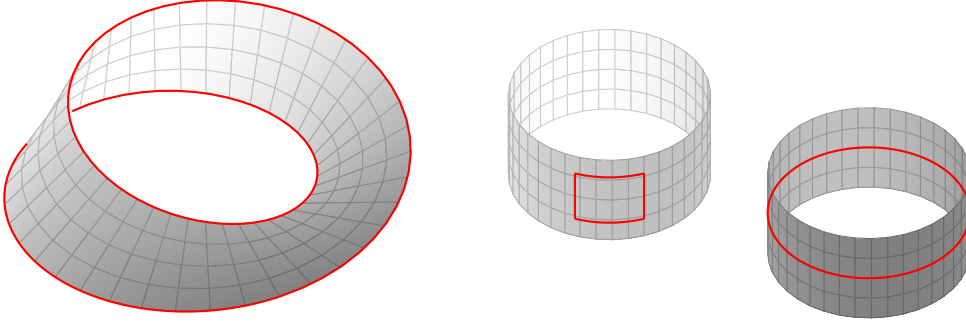


Figura 2.3. Curvas cerradas en el cilindro.

2.8 Ejemplo. Tomando $\mathbb{R}^k \hookrightarrow \mathbb{R}^n$, para $k < n$, se tiene que $\mathbb{R}^n \setminus \mathbb{R}^k \equiv \mathbb{S}^{n-k-1}$. Haremos una prueba inductiva sobre todas las parejas de la forma (n, k) con $k < n$. Para el caso $(n, 1)$ consideremos:

$$\begin{aligned} f : \mathbb{R}^n \setminus \mathbb{R}^1 &\rightarrow \mathbb{S}^{n-2}, \\ (x_1, \dots, x_n) &\mapsto \left(\frac{x_2}{\sqrt{x_2^2 + \dots + x_n^2}}, \dots, \frac{x_n}{\sqrt{x_2^2 + \dots + x_n^2}} \right), \\ g : \mathbb{S}^{n-2} &\rightarrow \mathbb{R}^n \setminus \mathbb{R}^1, \\ (x_1, \dots, x_{n-1}) &\mapsto (0, x_1, \dots, x_{n-1}), \end{aligned}$$

De esto obtendremos $f \circ g = \text{Id}_{\mathbb{S}^{n-2}}$ y además

$$(g \circ f)(x_1, \dots, x_n) = \left(0, \frac{x_2}{\sqrt{x_2^2 + \dots + x_n^2}}, \dots, \frac{x_n}{\sqrt{x_2^2 + \dots + x_n^2}} \right),$$

y definiendo, para $x = (x_1, \dots, x_n) \in \mathbb{R}^n \setminus \mathbb{R}^1$:

$$H(x, t) = \left(tx_1, tx_2 + (1-t) \frac{x_2}{\sqrt{x_2^2 + \dots + x_n^2}}, \dots, tx_n + (1-t) \frac{x_n}{\sqrt{x_2^2 + \dots + x_n^2}} \right),$$

Obtendremos una homotopía entre $g \circ f$ y $\text{Id}_{\mathbb{R}^n \setminus \mathbb{R}^1}$, de esto que se verifica el caso $(n, 1)$ para cada n . Nuestra estrategia ahora será probar que si asumimos el caso (n, k) esto implica el caso $(n+1, k+1)$ y de esta manera estaremos llenando todos

los casos que nos interesan. Para ello, podemos ver que $\mathbb{R}^{n+1} \setminus \mathbb{R}^{k+1}$ es un espacio homeomorfo a $\mathbb{R} \times (\mathbb{R}^n \setminus \mathbb{R}^k)$ y por tanto este tendrá el mismo tipo de homotopía que $(\mathbb{R}^n \setminus \mathbb{R}^k)$ pues $\mathbb{R}$ se puede retraer a un punto. Obteniendo así las siguientes equivalencias homotópicas:

$$\mathbb{R}^{n+1} \setminus \mathbb{R}^{k+1} \equiv \mathbb{R} \times (\mathbb{R}^n \setminus \mathbb{R}^k) \equiv (\mathbb{R}^n \setminus \mathbb{R}^k) \equiv \mathbb{S}^{n-k-1},$$

y claramente $\mathbb{S}^{n-k-1} = \mathbb{S}^{(n+1)-(k+1)-1}$. De esta manera, estaremos llenando (de manera similar que la inducción) todos puntos $(n, k) \in \mathbb{Z}^+ \times \mathbb{Z}^+$ tales que $n > k$.

Notemos que este resultado sigue siendo válido aún cuando $k = 0$, considerando a $\mathbb{R}^0$ como el origen de $\mathbb{R}^n$. En este caso, es un buen ejercicio para familiarizarse con las definiciones probar que $\mathbb{R}^n \setminus \mathbb{R}^0$ tiene el mismo tipo de homotopía que $\mathbb{S}^{n-1}$ para todo n .

2.9 Ejemplo. Los grupos $GL^+(n, \mathbb{C})$ (matrices invertibles con determinante positivo) y $SU(n, \mathbb{C})$ (matrices unitarias con determinante uno) tienen el mismo tipo de homotopía cuando se considera la topología usual inducida por cualquier norma. Para ver esto, consideremos la descomposición polar de $A \in GL^+(n, \mathbb{C})$, para toda matriz invertible A , esta puede ser escrita de manera única como:

$$A = UP,$$

donde U es una matriz unitaria y $P = \sqrt{A^*A}$. Esta unicidad nos permite construir la función $f : GL^+(n, \mathbb{C}) \rightarrow SU(n, \mathbb{C})$ definida como $f(A) = A \left(\sqrt{A^*A} \right)^{-1}$, donde la imagen de f es concretamente $SU(n, \mathbb{C})$ pues:

$$\overline{\det(A)} \det(A) = \det(A^*) \det(A) = \det(A^*A) = \left[\det \left(\sqrt{A^*A} \right) \right]^2,$$

sin embargo, $\det(A) > 0$ por lo que $\overline{\det(A)} = \det(A)$. Luego se tiene la igualdad $\left[\det \left(\sqrt{A^*A} \right) \right]^2 = [\det(A)]^2$ y de esto se deduce que $\det \left(\sqrt{A^*A} \right) = [\det(A)]$. Finalmente para $A \in GL^+(n, \mathbb{C})$ arbitraria se verifica:

$$\det(f(A)) = \det(A) \cdot \left[\det \left(\sqrt{A^*A} \right) \right]^{-1} = \det(A) \cdot [\det(A)]^{-1} = 1,$$

por lo que podemos considerar, abusando de la notación, $f : GL^+(n, \mathbb{C}) \rightarrow SU(n, \mathbb{C})$. Luego f y la inclusión $\iota : SU(n, \mathbb{C}) \hookrightarrow GL^+(n, \mathbb{C})$ son inversos homotópicos pues

$f \circ \iota = \text{Id}_{SU(n, \mathbb{C})}$ y $\iota \circ f \simeq \text{Id}_{GL^+(n, \mathbb{C})}$ ya que podemos considerar la homotopía lineal:

$$H(A, t) = (1 - t)A + tA \left(\sqrt{A^* A} \right)^{-1}.$$

El que esta homotopía tome valores en $GL^+(n, \mathbb{C})$ viene de que puede escribirse como :

$$A \left((1 - t)I + t \left(\sqrt{A^* A} \right)^{-1} \right),$$

donde $(1 - t)I + t \left(\sqrt{A^* A} \right)^{-1}$ es una matriz definida positiva pues es una combinación convexa de matrices definidas positivas y al calcular $\det H(A, t)$ este da como resultado el producto de dos números positivos. La continuidad de $H(A, t)$ viene de esta función está definida mediante composiciones de operaciones continuas entre matrices: multiplicar por escalares, producto de matrices, tomar traspuesta conjugada, la inversión y la raíz cuadrada. Para esta última el lector puede consultar [30, p. 411] para una construcción de la prueba.

2.2. Grupo fundamental

El grupo fundamental es una de las herramientas más intuitivas dentro de la topología algebraica y uno de los primeros invariantes que se consideran al momento de trabajar con variedades. El primero en estudiarlo fue Henri Poincaré a finales del siglo diecinueve. Antes de entrar de lleno con el grupo fundamental, exploraremos una construcción un tanto más general y se hará mención de algunas de sus propiedades a manera de estudio complementario.

2.2.1. Grupoide fundamental

Recordemos que a cada función continua $a : I \rightarrow X$, donde X es un espacio topológico, es llamada *camino* en X . El estudio de los posibles caminos en un espacio topológicos nos da una noción de conexidad más intuitiva y más fuerte que la conjuntista, por ende resulta natural tratar de investigar las clases de equivalencias de estos caminos mediante homotopías. En concreto, una homotopía de caminos será una homotopía estacionaria con $A = \partial I$, más concretamente:

2.10 Definición. Si $a, b : I \rightarrow X$, estos se dicen *caminos homotópicos* si existe una

aplicación continua $H : I \times I \rightarrow X$ tal que:

$$\begin{aligned} H(s, 0) &= a(s), H(s, 1) = b(s), \\ H(0, t) &= a(0) = b(0), \\ H(1, t) &= a(1) = b(1), \end{aligned}$$

para todo $s, t \in I$, o en otras palabras $a \simeq_{\partial I} b$.

Es importante diferenciar la imagen un camino del camino como tal, por ejemplo si a, b son caminos donde $a(s) = e^{2\pi is}$ y $b(s) = e^{4\pi is}$, estos son dos caminos distintos con la misma imagen.

2.11 Definición. Sean $a, b : I \rightarrow X$ caminos tales que $a(1) = b(0)$. Se define el *producto de caminos* $a \cdot b$ como el camino que consiste en recorrer primero a y luego b . Más precisamente:

$$a \cdot b(s) = \begin{cases} a(2s) & \text{si } 0 \leq s \leq 1/2 \\ b(2s - 1) & \text{si } 1/2 \leq s \leq 1. \end{cases}$$

Verificar que $a \cdot b : I \rightarrow X$ es continua, es una aplicación del Lema del pegado.

2.12 Definición. Sea $a : I \rightarrow X$ un camino. Se define el *reverso* $\bar{a}$ como el camino que consiste en recorrer a al revés. Más precisamente:

$$\bar{a}(s) = a(1 - s).$$

Así mismo, al *camino constante* que cuya imagen es el punto x lo denotamos por e_x , es decir $e_x : I \rightarrow X$ es tal que $e_x(s) = x$ para todo $s \in I$.

Con estas definiciones podríamos comenzar a pensar en una estructura algebraica que involucre caminos, sin embargo aún está un poco lejos de ser interesante. Esto se debe a que no siempre tendremos asociatividad y eso algo que buscamos en la mayoría de estructuras. Por ello veamos que al considerar las relaciones de equivalencia generada por la relación de homotopía de caminos se ganarán algunas propiedades. Si a, a', b, b' son caminos con $a \simeq_{\partial I} a'$ y $b \simeq_{\partial I} b'$, entonces $a \cdot b \simeq_{\partial I} a' \cdot b'$, siempre que los productos de caminos estén definidos. Esto puede verificarse definiendo:

$$H(s, t) = \begin{cases} F(2s, t) & \text{si } 0 \leq s \leq 1/2 \\ G(2s - 1, t) & \text{si } 1/2 \leq s \leq 1, \end{cases}$$

donde F es una homotopía de caminos entre a y a' , y G es homotopía de caminos entre b y b' . Estas relaciones nos indican que si $[a]$ es la clase de homotopía de caminos de a el producto entre clases

$$[a][b] := [a \cdot b],$$

esta bien definido y tiene las siguientes propiedades:

2.13 Proposición. *Si $a, b, c : I \rightarrow X$ son caminos en un espacio topológico X con $a(0) = x_0$ y $a(1) = x_1$ entonces*

1. *(Asociatividad) Si $[a]([b][c])$ está bien definido, también lo está $([a][b])[c]$ y estos son iguales.*
2. *(Identidades) Si $\varepsilon_x = [e_x]$ entonces $[a]\varepsilon_{x_1} = [a] = \varepsilon_{x_0}[a]$.*
3. *(Inversos) Si $[a]^{-1} = [\bar{a}]$ se tiene $[a]^{-1}[a] = \varepsilon_{x_1}$ y $[a][a]^{-1} = \varepsilon_{x_0}$.*

Aunque es una proposición muy común en libros de topología, presentaremos la demostración de [26, pp. 326-329] pues esta desarrolla ideas bastante generales que pueden ser usados para probar otras proposiciones.

Demostración. La prueba se basa en utilizar dos hechos sobre las homotopías:

- Si $k : X \rightarrow Y$ es una función continua, y F es una homotopía de caminos en X entre a y a' , entonces $k \circ F$ es una homotopía de caminos entre los caminos $k \circ a$ y $k \circ a'$.
- Si $k : X \rightarrow Y$ es una función continua, y si a y b son caminos en X tales que $a(1) = b(0)$, entonces

$$k \circ (a \cdot b) = (k \circ a) \cdot (k \circ b).$$

Se verificarán las propiedades 2 y 3. Para la 2, sea e_0 el camino constante de I en 0, notemos que Id_I puede ser visto como un camino de 0 a 1 en I . De esta forma $e_0 \cdot \text{Id}_I$ es también un camino en I de 0 a 1. Por la convexidad de I es posible construir una homotopía de caminos G entre Id_I y $e_0 \cdot \text{Id}_I$. Entonces $a \circ G$ es una homotopía de caminos en X entre los caminos $a \circ \text{Id}_I = a$ y

$$a \circ (e_0 \cdot \text{Id}_I) = (a \circ e_0) \cdot (a \circ \text{Id}_I) = e_{x_0} \cdot a,$$

esto prueba que

$$\varepsilon_{x_0}[a] = [e_{x_0}][a] = [e_{x_0} \cdot a] = [a],$$

y un argumento similar puede construirse para mostrar que $[a]\varepsilon_{x_1} = [a]$, notando primero que $\text{Id}_I \cdot e_1$ es un camino homotópico a Id_I y procediendo de manera análoga. Para probar β , notemos que $\overline{\text{Id}_I}(s) = 1 - s$. Entonces $\text{Id}_I \cdot \overline{\text{Id}_I}$ es un camino en I empezando y terminando en 0, tal como el camino constante e_0 . Como I es convexo, existe una homotopía de caminos H entre e_0 y $\text{Id}_I \cdot \overline{\text{Id}_I}$. Entonces $a \circ H$ es una homotopía de caminos entre $a \circ e_0 = e_{x_0}$ y

$$(a \circ \text{Id}_I) \cdot (a \circ \overline{\text{Id}_I}) = a \cdot \bar{a},$$

y un argumento similar sobre $\overline{\text{Id}_I} \cdot \text{Id}_I$ y e_{x_1} muestra $[a]^{-1}[a] = \varepsilon_{x_1}$.

Probar 1 tomará un poco más de trabajo. Para esta parte conviene describir el producto $a \cdot b$ de otra manera. Si $[x, y]$ y $[w, z]$ son dos intervalos en $\mathbb{R}$, existe una única aplicación de la forma $p(t) = mt + k$ que lleva y a z y x a w ; la cual es llamada *mapeo lineal positivo* de $[x, y]$ a $[w, z]$. Nótese que la inversa de esta función es otra de estas funciones. Con esta terminología tendremos que $a \cdot b$ puede ser descrito de la siguiente forma: en $[0, 1/2]$, es igual al mapeo lineal positivo de $[0, 1/2]$ a $[0, 1]$ compuesto con a y en $[1/2, 1]$ es el mapeo lineal positivo de $[1/2, 1]$ a $[0, 1]$ compuesto con b . Ahora verificaremos 1. Dados los caminos a, b y c en X , los productos $a \cdot (b \cdot c)$ y $(a \cdot b) \cdot c$ están definidos precisamente cuando $a(1) = b(0)$ y $b(1) = c(0)$. Asumiendo estas dos condiciones, definimos también el “triple producto” de los caminos a, b, c como: se escogen los puntos x, y de I tales que $0 < x < y < 1$. Se define el camino $k_{x,y}$ como: en $[0, x]$ el mapeo lineal positivo de $[0, x]$ a I compuesto con a , en $[x, y]$ el mapeo lineal positivo de $[x, y]$ a I compuesto por b ; y en $[y, 1]$ el mapeo lineal positivo de $[y, 1]$ a I compuesto con c . Es posible mostrar usando la convexidad de I que para cualquier otra elección de $z, w \in I$ tal que $0 < w < z < 1$ que existe una aplicación $p : I \rightarrow I$ de tal forma que $k_{w,z} \circ p = k_{x,y}$, donde esta es homotópica a la identidad, esto muestra que $k_{x,y}$ y $k_{w,z}$ son caminos homotópicos. Finalmente, podemos notar que $a \cdot (b \cdot c) = k_{\frac{1}{2}, \frac{3}{4}} \simeq_{\partial I} k_{\frac{1}{4}, \frac{1}{2}} = (a \cdot b) \cdot c$. ■

Aunque un poco extensa, esta idea en la demostración da lugar a otra proposición bastante útil y es por eso que se opta por esta prueba sobre otras presentes en la literatura. Dicha proposición es la siguiente:

2.14 Proposición. Sea f un camino en X , y sean $a_0, \dots, a_n$ números tales que $0 = a_0 < a_1 < \dots < a_n = 1$. Sea f_i el mapeo lineal positivo sobreyectivo de I a $[a_{i-1}, a_i]$ compuesto con f . Entonces

$$[f] = [f_1][f_2] \cdots [f_n].$$

A la colección de todas las clases de equivalencia en un espacio topológico X se denota por $\Pi(X)$ y se le conoce como *grupoide fundamental*. En general un grupoide puede verse como una categoría que satisface estas tres propiedades de la proposición anterior. Así mismo es posible definir un isomorfismo para grupoides y mostrar que, en realidad, el grupoide fundamental es un invariante algebraico, esto significa que si X y Y son dos espacios topológicos homeomorfos, entonces $\Pi(X)$ y $\Pi(Y)$ son dos grupoides isomorfos. Si se desea profundizar en el estudio de grupoides puede consultarse el capítulo 6 de [5], en donde se profundiza en la idea de grupoides como categorías. En particular, se recomienda esta lectura pues da una introducción moderna a la topología general y profundiza en temas poco comunes en la materia.

2.2.2. Definición y propiedades de π_1 .

Aunque la estructura de grupoide es más general que la de un grupo, por ahora fijaremos nuestra atención en un grupo muy particular que puede ser extraído de un espacio topológico.

2.15 Definición. Un *lazo* en X , un espacio topológico, es un camino $a : I \rightarrow X$ tal que $a(0) = a(1)$. Si $p = a(0) = a(1)$ diremos que a está *basado en p* o que a *tiene como base el punto p* . Al conjunto de todos los lazos en un espacio topológico X con base en un punto p lo denotaremos por $\Omega(X, p)$, esto es:

$$\Omega(X, p) = \{a \in C(I, X) : p = a(0) = a(1)\}.$$

Al restringirnos a un punto en concreto las propiedades de la Proposición 2.13 nos aseguran que $\Omega(X, p) / \simeq_{\partial I}$ es un grupo donde la identidad es justo la clase de el camino constante en p , a este grupo lo llamaremos el *grupo fundamental de X con base en p* y se denota por $\pi_1(X, p)$.

Notemos que este grupo depende de la elección de nuestro punto base, sin embargo, muchas veces esto puede omitirse cuando se trabaja con espacios conexos por caminos, esto se debe al siguiente teorema:

2.16 Teorema. Sea X conexo por caminos, si $p, q \in X$ y b es cualquier camino de p a q . El mapa $\Phi_b : \pi_1(X, p) \rightarrow \pi_1(X, q)$ definido por:

$$\Phi_b([a]) = [\bar{b}][a][b],$$

es un isomorfismo de grupos.

Demostración. Notemos que $\bar{b}$ es un camino de q a p y a es un lazo, es decir, un camino de p a p , por lo que los productos de Φ_b siempre estarán definidos y $\bar{b} \cdot a \cdot b$ es un lazo basado en q . Ahora, en efecto es un homomorfismo pues, por la Proposición 2.13 se tiene:

$$\begin{aligned} \Phi_b([a_1])\Phi_b([a_2]) &= [\bar{b}][a_1][b][\bar{b}][a_2][b] \\ &= [\bar{b}][a_1][c_p][a_2][b] \\ &= [\bar{b}][a_1][a_2][b] \\ &= \Phi_b([a_1][a_2]). \end{aligned}$$

Y repitiendo el mismo procedimiento para $\Phi_{\bar{b}}$ puede verificarse que este es un homomorfismo de $\pi_1(X, q)$ a $\pi_1(X, p)$, siendo este el inverso de Φ_b . ■

Por ello cuando tengamos un espacio conexo por caminos a veces se omite el punto base de los lazos y se escribe simplemente $\pi_1(X)$. Algunos de los resultados importantes relacionados al grupo fundamental son los siguientes:

2.17 Proposición. Si $\varphi : X \rightarrow Y$ es continua, la aplicación $\varphi_* : \pi_1(X, p) \rightarrow \pi_1(Y, \varphi(p))$ definida por:

$$\varphi_*([a]) = [\varphi \circ a],$$

es un homomorfismo. A φ_* se le conoce como homomorfismo inducido por φ .

Demostración. Primero, está bien definido pues, si $[a_1] = [a_2]$, entonces $[\varphi \circ a_1] = [\varphi \circ a_2]$ pues es uno de los hechos sobre homotopías en la prueba de la Proposición 2.13. La condición $\varphi_*([a][b]) = \varphi_*([a])\varphi_*([b])$ viene del segundo hecho sobre homotopías en la prueba de 2.13. ■

En el caso particular en que φ es un homeomorfismo:

2.18 Proposición. Si $\varphi : X \rightarrow Y$ es un homeomorfismo, φ_* es un isomorfismo de grupos. Esto es, si X, Y son dos espacios homeomorfos con $p \in X$, entonces $\pi_1(X, p)$ y $\pi_1(Y, \varphi(p))$ son isomorfos.

Una de las maneras más prácticas para calcular un grupo fundamental es mostrar que el espacio del cual queremos calcular el grupo fundamental es homeomorfo a otro del cual tenemos más información o, de manera similar, usar la siguiente proposición:

2.19 Proposición. Si $\varphi : X \rightarrow Y$ es una equivalencia homotópica, entonces para cada $p \in X$, $\varphi_* : \pi_1(X, p) \rightarrow \pi_1(Y, \varphi(p))$ es un isomorfismo. En otras palabras si X, Y tienen el mismo tipo de homotopía, sus grupos fundamentales son isomorfos.

2.20 Lema (Lema del cuadrado). Sea $F : I \times I \rightarrow X$ una función continua, y sean a, b, c y d caminos en X definidos por:

$$\begin{aligned} a(s) &= F(s, 0), & b(s) &= F(1, s), \\ c(s) &= F(0, s), & d(s) &= F(s, 1). \end{aligned}$$

Entonces $a \cdot b \simeq c \cdot d$.

Demostración. Se considera la homotopía

$$H(s, t) = \begin{cases} F(2st, 2s(1-t)) & \text{si } 0 \leq s \leq 1/2 \\ F((2s-1)(1-t) + t, (2s-1)t + (1-t)) & \text{si } 1/2 \leq s \leq 1 \end{cases}.$$

La cual verifica

$$\begin{aligned} H(s, 0) &= \begin{cases} F(0, 2s) & \text{si } 0 \leq s \leq 1/2 \\ F((2s-1), 1) & \text{si } 1/2 \leq s \leq 1 \end{cases} = (c \cdot d)(s); \\ H(s, 1) &= \begin{cases} F(2s, 0) & \text{si } 0 \leq s \leq 1/2 \\ F(1, 2s-1) & \text{si } 1/2 \leq s \leq 1 \end{cases} = (a \cdot b)(s). \end{aligned}$$

Veamos que H está bien definida cuando $s = 1/2$ y es una composición de funciones continuas en los conjuntos cerrados $[0, \frac{1}{2}] \times I$ y $[\frac{1}{2}, 1] \times I$, por el Lema del pegado (ver Apéndice A, Proposición A.18), H es continua en $[0, \frac{1}{2}] \times I \cup [\frac{1}{2}, 1] \times I = I \times I$, por ello, se concluye que $a \cdot b \simeq c \cdot d$. ■

2.21 Lema. Sean $\varphi, \psi : X \rightarrow Y$ continuas y $H : \varphi \simeq \psi$ una homotopía. Para cada $p \in X$, sea h el camino en Y de $\varphi(p)$ a $\psi(p)$ definido por $h(t) = H(p, t)$, y

sea $\Phi_h : \pi_1(Y, \varphi(p)) \rightarrow \pi_1(Y, \psi(p))$ el isomorfismo $\Phi_h([a]) = [\bar{h}][a][h]$. Entonces, el siguiente diagrama conmuta

$$\begin{array}{ccc}
 & \pi_1(Y, \varphi(p)) & \\
 \varphi_* \nearrow & \downarrow \Phi_h & \\
 \pi_1(X, p) & & \pi_1(Y, \psi(p)). \\
 \psi_* \searrow & &
 \end{array}$$

Demostración. Sea f un lazo en X con base en el punto p . Lo que queremos mostrar es $\psi_*[f] = \Phi_h(\varphi_*[f])$, lo cual es equivalente a $h \cdot (\psi \circ f) \simeq (\varphi \circ f) \cdot h$, siendo esto cierto por el lema del cuadrado tomando $F(s, t) = H(f(s), t)$, o visto geoméricamente en la Figura 2.4. ■

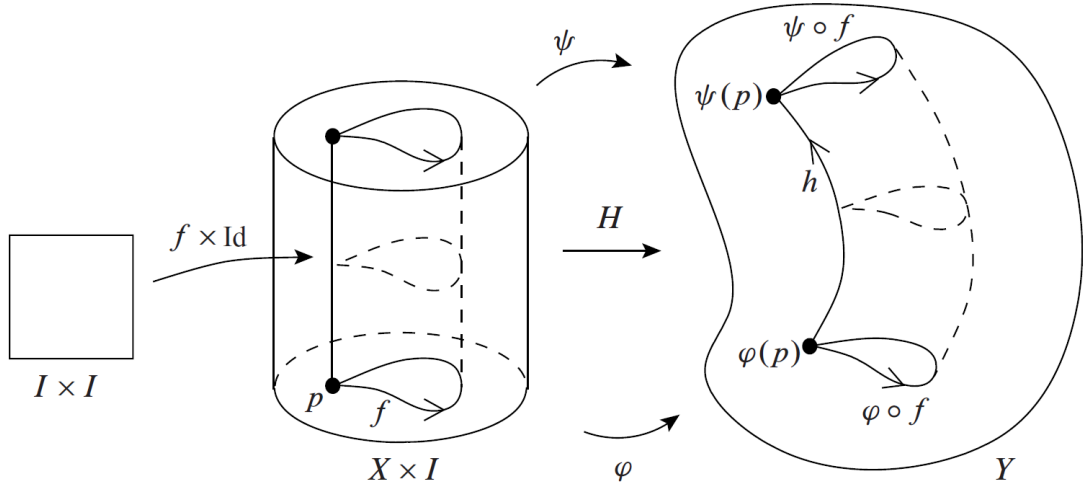


Figura 2.4. Homomorfismos inducidos por aplicaciones homotópicas. Fuente: tomada de [19, p. 204]

Demostración de 2.19. Sea $\varphi : X \rightarrow Y$ una equivalencia homotópica y sea $\psi : Y \rightarrow X$ el inverso homotópico de φ . Si se considera la siguiente sucesión de aplicaciones:

$$\pi_1(X, p) \xrightarrow{\varphi_*} \pi_1(Y, \varphi(p)) \xrightarrow{\psi_*} \pi_1(X, \psi(\varphi(p))) \xrightarrow{\varphi_*} \pi_1(Y, \varphi(\psi(\varphi(p)))).$$

Hay que probar que φ_* es biyectivo. Como pudimos ver anteriormente, ψ_* no es una inversa pues no siempre se tiene como imagen el grupo que buscamos. Como

$\psi \circ \varphi \simeq \text{Id}_X$, el Lema 2.21 muestra que existe un camino h en X tal que el siguiente diagrama conmuta:

$$\begin{array}{ccc}
 & \pi_1(X, p) & \\
 (\text{Id}_X)_* \nearrow & \downarrow \Phi_h & \\
 \pi_1(X, p) & & \pi_1(X, \psi(\varphi(p))). \\
 (\psi \circ \varphi)_* \searrow & &
 \end{array}$$

Por lo que $\psi_* \circ \varphi_* = \Phi_h$, es un isomorfismo. En particular muestra que φ_* es inyectivo y ψ_* es sobreyectivo. De manera similar, para $\varphi \circ \psi \simeq \text{Id}_Y$ se tiene el diagrama

$$\begin{array}{ccc}
 & \pi_1(Y, \varphi(p)) & \\
 (\text{Id}_Y)_* \nearrow & \downarrow \Phi_k & \\
 \pi_1(Y, \varphi(p)) & & \pi_1(Y, \varphi(\psi(\varphi(p)))). \\
 (\varphi \circ \psi)_* \searrow & &
 \end{array}$$

Que por un argumento análogo muestra que ψ_* inyectivo, y por tanto biyectivo. De esto que $\varphi_* = (\psi_*)^{-1} : \pi_1(X, p) \rightarrow \pi_1(Y, \varphi(p))$ es un isomorfismo. ■

En algunos espacios no podemos dar explícitamente una equivalencia homotópica, sin embargo, es posible al menos dar una relación entre dos grupos fundamentales.

2.22 Proposición. Si $r : X \rightarrow A$ es una función continua con $A \subseteq X$ que verifica $r|_A = \text{Id}_A$. Entonces, para cada $p \in A$, se tiene que $(\iota_A)_* : \pi_1(A, p) \rightarrow \pi_1(X, p)$ es un homomorfismo inyectivo, donde ι_A es la aplicación inclusión de A en X y $r_* : \pi_1(X, p) \rightarrow \pi_1(A, p)$ es un homomorfismo sobreyectivo.

En general a una función continua $r : X \rightarrow A$ que verifique $r|_A = \text{Id}_A$ o equivalentemente $r \circ \iota_A = \text{Id}_A$, es llamada *retracción*, además, si $\iota_A \circ r$ es homotópica a Id_X , r es una *deformación retráctil*.

2.23 Ejemplo. Un *conjunto en forma de estrella* es un subconjunto S de $\mathbb{R}^n$ para el que existe $p_0 \in S$ tal que para todo $p \in S$ el segmento de recta que une a p_0 y p está contenido en S . Más precisamente si $p \in S$ entonces $\lambda p_0 + (1 - \lambda)p \in S$ para todo $\lambda \in I$. Estos conjuntos tienen grupo fundamental trivial pues existe una deformación retráctil de el conjunto entero al punto p_0 de nuestra definición, como

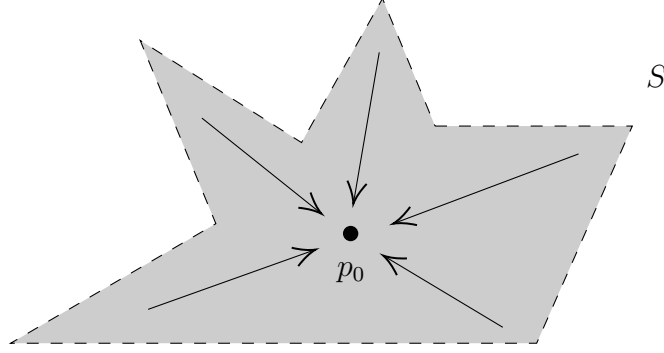


Figura 2.5. Deformación retráctil de un conjunto en forma de estrella.

podemos ver en la Figura 2.5. En concreto los conjuntos convexos son conjuntos en forma de estrella.

2.24 Ejemplo. Para las esferas $\mathbb{S}^n \subset \mathbb{R}^{n+1}$ con $n \geq 1$ se verifica:

$$\pi_1(\mathbb{S}^n) \cong \begin{cases} \mathbb{Z} & n = 1 \\ \{e\} & n \geq 2 \end{cases}$$

Para pruebas detalladas de esto, pueden consultarse [21], [28] y [26].

Finalmente, un resultado que involucra un producto cartesiano de dos espacios topológicos y facilita la tarea de calcular grupos fundamentales, es la siguiente

2.25 Proposición. Si $X_1, \dots, X_n$ son espacios topológicos con $x_i \in X_i$ para cada $1 \leq i \leq n$, se tiene el siguiente isomorfismo de grupos:

$$\pi_1(X_1 \times \dots \times X_n, (x_1, \dots, x_n)) \cong \pi_1(X_1, x_1) \times \dots \times \pi_1(X_n, x_n).$$

2.2.3. Espacios simplemente conexos

Usualmente, en un primer curso de cálculo de varias variables o un curso de variable compleja se habla subconjuntos de $\mathbb{R}^2$ como simplemente conexos si “no tienen agujeros”, sin embargo, una definición mucho más precisa y general (incluso para conjuntos que no estén contenidos en $\mathbb{R}^2$) puede darse en función del grupo fundamental.

2.26 Definición. Sea X un espacio topológico, este se dice *simplemente conexo* si X es conexo por caminos y $\pi_1(X)$ es un grupo trivial.

Esta definición es bastante general, pues ni siquiera hace referencia a la topología de algún $\mathbb{R}^n$. A pesar de eso, algunos autores como Ahlfors en [2] han optado

(por fines prácticos) en dar definiciones menos generales, en concreto la definición dada por Ahlfors es la siguiente:

2.27 Definición. Una subconjunto no vacío, abierto y conexo de $\mathbb{C}$ es llamado simplemente conexo si su complemento respecto al plano extendido $\mathbb{C} \cup \{\infty\}$ es conexo.

Estas dos definiciones son equivalentes (por lo menos en $\mathbb{C}$) y muchas de las proposiciones en el análisis de variable compleja, en su forma más general, involucran este tipo de conjuntos. A continuación se presentan algunos ejemplos de conjuntos simplemente conexos.

2.28 Ejemplo. Todos los conjuntos X que pueden contraerse a un punto son simplemente conexos, es decir, todos para los que existe una deformación retráctil $r : X \rightarrow \{p\}$ son simplemente conexos. A estos se les conoce como espacios *contráctiles* o *contraíbles*. Más específicamente, los conjuntos en forma de estrella son simplemente conexos.

2.29 Ejemplo. Sea V un espacio vectorial sobre $\mathbb{R}$ de dimensión n con la topología usual generada por cualquier norma y $U \subset V$ es un subespacio cuya dimensión es $k < n - 2$, (es decir, un espacio de codimensión mayor a 2) la diferencia $V \setminus U$ es un conjunto simplemente conexo. Para mostrar esto, consideremos una base de U dada por $\{u_1, \dots, u_k\}$, si se extiende a una base $\{u_1, \dots, u_k, v_{k+1}, \dots, v_n\}$ de V . Se define la siguiente transformación lineal:

$$T(a_1u_1 + \dots + a_ku_k + a_{k+1}v_{k+1}, \dots, a_nv_n) = (a_1, \dots, a_n).$$

Como toda transformación lineal entre espacios de dimensión finita es continua, se tiene que T es continua y al ser invertible también esta inversa es continua, obteniendo un homeomorfismo entre V y $\mathbb{R}^n$ tal que $T(U) = \mathbb{R}^k \hookrightarrow \mathbb{R}^n$. De esta manera, como $T(V \setminus U) = T(V) \setminus T(U) = \mathbb{R}^n \setminus \mathbb{R}^k$ obtendremos que $V \setminus U$ es homeomorfo a $\mathbb{R}^n \setminus \mathbb{R}^k$. Por el Ejemplo 2.8, $V \setminus U$ tiene el mismo tipo de homotopía que $\mathbb{S}^{n-k-1}$ y por tanto se tendrá que $\pi_1(V \setminus U) \cong \pi_1(\mathbb{S}^{n-k-1})$ donde $n - k - 1 \geq 2$ y por tanto $\pi_1(V \setminus U)$ es simplemente conexo.

2.3. Presentaciones de grupos

Al igual que con espacios topológicos, podemos crear más de un espacio topológico tomando como base dos espacios y alguna operación entre ellos, resulta algo muy parecido con grupos. En este caso exploraremos otra forma de crear un producto de grupos distinto del producto directo o del producto semidirecto de dos grupos. Para ello primero exploraremos un poco de teoría combinatoria de grupos.

2.30 Definición. Dada una familia indizada de grupos $(G_\alpha)_{\alpha \in \Lambda}$ una *palabra* es una m -tupla con $m \geq 0$ de elementos en la unión disjunta $G(\Lambda) = \bigsqcup_{\alpha \in \Lambda} G_\alpha$. Más concretamente una palabra es un elemento de

$$\bigcup_{n=1}^{\infty} \underbrace{G(\Lambda) \times \cdots \times G(\Lambda)}_n,$$

o la tupla vacía $()$. Al conjunto de todas las palabras en $(G_\alpha)_{\alpha \in \Lambda}$ lo llamaremos $\mathcal{W}$.

Notemos que los elementos de la unión disjunta son de la forma (g, α) , donde $g \in G_\alpha$ para algún $\alpha \in \Lambda$, sin embargo, para abreviarlo a veces se opta por usar la notación g_α . De esta manera, los elementos de $\mathcal{W}$ serán de la forma $(g_1, \dots, g_m)$. Al elemento neutro de cada G_α lo llamaremos 1_α .

2.31 Definición. Si $(g_1, \dots, g_m), (h_1, \dots, h_k) \in \mathcal{W}$ se define el producto de palabras como la siguiente yuxtaposición:

$$(g_1, \dots, g_m)(h_1, \dots, h_k) = (g_1, \dots, g_m, h_1, \dots, h_k).$$

Notemos que este producto es asociativo y la palabra vacía $()$ es un neutro por ambos lados, sin embargo, aún carecemos de inversos para llamar a nuestra estructura grupo, además, realmente no estamos aprovechando las propiedades de cada G_α como grupo. Para arreglar estos problemas consideraremos la siguiente relación de equivalencia:

2.32 Definición. Una *reducción elemental* es una de las siguientes operaciones:

$$\begin{aligned} (g_1, \dots, g_i, g_{i+1}, \dots, g_m) &\mapsto (g_1, \dots, g_i g_{i+1}, \dots, g_m) \quad \text{si } g_i, g_{i+1} \in G_\alpha, \\ (g_1, \dots, g_{i-1}, 1_\alpha, g_{i+1}, \dots, g_m) &\mapsto (g_1, \dots, g_{i-1}, g_{i+1}, \dots, g_m). \end{aligned}$$

Ahora, se toma la relación de equivalencia $\sim$ en $\mathcal{W}$ definida de la siguiente forma: sean $W, W' \in \mathcal{W}$ entonces $W \sim W'$ si y solo si existe una sucesión finita de

palabras $W = W_0, W_1, \dots, W_n = W'$ tales que para cada $i = 1, \dots, n$, la palabra W_i puede ser obtenida de W_{i-1} por reducciones elementales o viceversa.

2.33 Definición. El conjunto cociente $\mathcal{W}/\sim$ es llamado el *producto libre* de los grupos $(G_\alpha)_{\alpha \in \Lambda}$, y se denota por $\bigstar_{\alpha \in \Lambda} G_\alpha$.

En caso que nuestro conjunto de índices Λ sea finito escribimos $G_1 * \dots * G_n$. Lo que procede a partir de ahora es mostrar que $\bigstar_{\alpha \in \Lambda} G_\alpha$ es un grupo.

2.34 Proposición. *Dada una familia indizada de grupos $(G_\alpha)_{\alpha \in \Lambda}$ el producto libre $\bigstar_{\alpha \in \Lambda} G_\alpha$ es un grupo respecto a la operación entre clases de equivalencia:*

$$[(g_1, \dots, g_m)] [(h_1, \dots, h_k)] = [(g_1, \dots, g_m, h_1, \dots, h_k)].$$

Demostración. Primero, supongamos que $W \sim W'$ y $V \sim V'$, un argumento inductivo sobre el número de reducciones elementales puede mostrar que $WV' \sim W'V'$ y $WV \sim WV'$ lo cual concluye en $[WV] = [W'V']$. Con la operación bien definida podemos ver que:

$$\begin{aligned} (g_1, \dots, g_m)(g_m^{-1}, \dots, g_1^{-1}) &\sim (), \\ (g_m^{-1}, \dots, g_1^{-1})(g_1, \dots, g_m) &\sim (). \end{aligned}$$

Esto muestra la existencia de inversos, teniendo así un grupo con neutro $[()]$. ■

Resulta que es conveniente elegir un solo representante para las clases de equivalencia. La existencia de este representante canónico viene por la siguiente proposición:

2.35 Proposición. *Se dice que una palabra es reducida, si no puede ser acortada mediante reducciones elementales. Todo elemento en $\bigstar_{\alpha \in \Lambda} G_\alpha$ es representado por una y solo una palabra reducida.*

La idea para la prueba de esta proposición consiste en dar un algoritmo explícito de como reducir una palabra y mostrar que, si $W \sim W'$ entonces al aplicar el algoritmo tanto en W como W' se obtiene la misma palabra. Para ver dicho algoritmo puede consultarse [28, pp. 32-34] y [19, pp. 235-237]. Debido a este resultado, a manera de convención se opta por referirnos a $[(g_1, \dots, g_m)]$ simplemente como $(g_1, \dots, g_m)$, o $g_1 \dots g_m$ cuando $(g_1, \dots, g_m)$ sea una palabra reducida.

Así mismo, podemos caracterizar el producto libre con la siguiente propiedad:

2.36 Teorema. Sea $(G_\alpha)_{\alpha \in \Lambda}$ una familia indizada de grupos. Para cualquier grupo H y toda colección de homomorfismos $\varphi_\alpha : G_\alpha \rightarrow H$, existe un único homomorfismo $\Phi : \bigstar_{\alpha \in \Lambda} G_\alpha \rightarrow H$ tal que para cada α el siguiente diagrama conmuta:

$$\begin{array}{ccc} \bigstar_{\alpha \in \Lambda} G_\alpha & & \\ \uparrow \iota_\alpha & \searrow \Phi & \\ G_\alpha & \xrightarrow{\varphi_\alpha} & H. \end{array}$$

Demostración. Abusaremos un poco de la notación identificando a G_α con su imagen bajo ι_α . Tomando los homomorfismos $\varphi_\alpha : G_\alpha \rightarrow H$, que el diagrama sea conmutativo y Φ sea un homomorfismo implica que se cumplen las siguientes condiciones:

$$\begin{aligned} \Phi(g) &= \varphi_\alpha(g) \quad \text{si } g \in G_\alpha, \\ \Phi(g_1 \cdots g_m) &= \Phi(g_1) \cdots \Phi(g_m) \end{aligned}$$

De esta manera si $g_i \in G_{\alpha_i}$ y $\bar{\Phi}$ es otro homomorfismo que verifica estas condiciones, entonces:

$$\bar{\Phi}(g_1 \cdots g_m) = \varphi_{\alpha_1}(g_1) \cdots \varphi_{\alpha_m}(g_m) = \Phi(g_1 \cdots g_m),$$

por lo que $\bar{\Phi} = \Phi$. Luego, para mostrar la existencia de dicho homomorfismo, para las palabras reducidas (las cuales por la proposición anterior sabemos que son únicas en cada clase de equivalencia) definimos la función directamente usando sus propiedades, solamente hay que mostrar que el morfismo está bien definido cuando se tienen reducciones elementales, para ello, tomando $g_1, g_2 \in G_\alpha$ para algún α , $g_1 g_2 \in G_\alpha \subseteq \bigstar_{\alpha \in \Lambda} G_\alpha$ es reducida o 1_α entonces, en el primer caso $\Phi(g_1 g_2) = \varphi_\alpha(g_1 g_2) = \varphi_\alpha(g_1) \varphi_\alpha(g_2) = \Phi(g_1) \Phi(g_2)$, por lo que sigue siendo un homomorfismo en este caso particular. En el segundo caso, por las propiedades de cada homomorfismo φ_α se tiene que $\Phi(1_\alpha) = \varphi_\alpha(1_\alpha) = 1 \in H$. Con esto mostrado, es posible construir una prueba inductiva para generalizar el argumento. ■

Esta propiedad puede ser usada para caracterizar el producto libre de grupos:

2.37 Proposición. El producto libre es el único grupo (salvo isomorfismo) que satisface la propiedad del Teorema 2.36.

Demostración. Supongamos que existe otro grupo $\mathcal{G}$ con esta propiedad, entonces, si $j_\alpha : G_\alpha \hookrightarrow \mathcal{G}$ para cada $\alpha \in \Lambda$ son las inclusiones, existen los homomorfismos

Φ_1 y Φ_2 tales que los siguientes diagramas conmutan:

$$\begin{array}{ccc} \bigstar_{\alpha \in \Lambda} G_\alpha & & \mathcal{G} \\ \uparrow \iota_\alpha & \searrow \Phi_1 & \uparrow j_\alpha \\ G_\alpha & \xrightarrow{j_\alpha} & \mathcal{G} \end{array} \quad \begin{array}{ccc} \mathcal{G} & & \bigstar_{\alpha \in \Lambda} G_\alpha \\ \uparrow j_\alpha & \searrow \Phi_2 & \uparrow \iota_\alpha \\ G_\alpha & \xrightarrow{j_\alpha} & \mathcal{G} \end{array}$$

Entonces se verifica $\Phi_2 \circ \Phi_1 \circ \iota_\alpha = \iota_\alpha$ y $\Phi_1 \circ \Phi_2 \circ j_\alpha = j_\alpha$ para cada $\alpha \in \Lambda$. Así que se tiene:

$$\begin{array}{ccc} \bigstar_{\alpha \in \Lambda} G_\alpha & & \mathcal{G} \\ \uparrow \iota_\alpha & \searrow \Phi_2 \circ \Phi_1 & \uparrow j_\alpha \\ G_\alpha & \xrightarrow{j_\alpha} & \mathcal{G} \end{array} \quad \begin{array}{ccc} \mathcal{G} & & \bigstar_{\alpha \in \Lambda} G_\alpha \\ \uparrow j_\alpha & \searrow \Phi_1 \circ \Phi_2 & \uparrow \iota_\alpha \\ G_\alpha & \xrightarrow{j_\alpha} & \mathcal{G} \end{array}$$

Sin embargo, las identidades $\text{Id}_{\bigstar_{\alpha \in \Lambda} G_\alpha} : \bigstar_{\alpha \in \Lambda} G_\alpha \rightarrow \bigstar_{\alpha \in \Lambda} G_\alpha$ y $\text{Id}_{\mathcal{G}} : \mathcal{G} \rightarrow \mathcal{G}$ también son homomorfismos que hacen conmutar los anteriores diagramas, por lo que, por la unicidad en la propiedad característica se tiene $\Phi_2 \circ \Phi_1 = \text{Id}_{\bigstar_{\alpha \in \Lambda} G_\alpha}$ y $\Phi_1 \circ \Phi_2 = \text{Id}_{\mathcal{G}}$. Es decir Φ_1 y Φ_2 , son isomorfismos de grupos. ■

2.38 Ejemplo. Si a es un generador de $\mathbb{Z}/3\mathbb{Z}$ y b es un generador de $\mathbb{Z}/2\mathbb{Z}$, los elementos de $\mathbb{Z}/3\mathbb{Z} * \mathbb{Z}/2\mathbb{Z}$ son todos de la forma $a^{m_1}b^{n_1}a^{m_2}b^{n_2} \dots a^{m_k}b^{n_k}$ donde $0 \leq m_i \leq 2$ y $0 \leq n_i \leq 1$ para $i = 1, \dots, k$. Este grupo no es abeliano pues

$$\begin{aligned} (aba^2b)(bab) &= a, \\ (bab)(aba^2b) &= bababa^2b. \end{aligned}$$

Por un lado tenemos una palabra reducida de longitud 1 y de otro, palabra reducida más grande, por lo que no pueden ser iguales. Esta es una de las diferencias respecto a un producto directo.

Otra herramienta importante dentro de la topología algebraica, es la construcción de un grupo dado un conjunto no vacío arbitrario S , sin que dicho conjunto tenga ninguna estructura. Algunas de estas construcciones son bastante generales e incluso es posible construir estructuras más robustas con ideas similares, sin embargo, por ahora nos conformaremos con obtener grupos.

2.39 Definición. Dado cualquier objeto σ , se define el grupo $F(\sigma)$, donde sus

elementos son $\{\sigma\} \times \mathbb{Z}$ y la operación binaria está dada por:

$$\begin{aligned} (\{\sigma\} \times \mathbb{Z}) \times (\{\sigma\} \times \mathbb{Z}) &\rightarrow (\{\sigma\} \times \mathbb{Z}), \\ ((\sigma, m), (\sigma, k)) &\mapsto (\sigma, m + k). \end{aligned}$$

Donde para simplificar notación escribimos σ^m en lugar de (σ, m) para cada $m \in \mathbb{Z}$. Al grupo $F(\sigma)$ lo llamaremos *grupo libre generado por σ* .

Este grupo puede verse como todas las potencias de σ con la multiplicación usual. Esencialmente $F(\sigma)$, es isomorfo a $(\mathbb{Z}, +)$, pero con un generador arbitrario σ .

2.40 Definición. Dado un conjunto arbitrario S . Se define el *grupo libre en S* como el producto libre:

$$F(S) = \bigstar_{\sigma \in S} F(\sigma).$$

Esta estructura también posee una propiedad que la caracteriza, la cual es la siguiente:

2.41 Teorema. Si S es un conjunto arbitrario. Para cada grupo H y cada mapeo $\varphi : S \rightarrow H$, existe un único homomorfismo $\Phi : F(S) \rightarrow H$ extendiendo φ :

$$\begin{array}{ccc} F(S) & & \\ \uparrow \iota & \searrow \Phi & \\ S & \xrightarrow{\varphi} & H. \end{array}$$

Demostración. Para cada función $\varphi : S \rightarrow H$ existen únicos homomorfismos $\varphi_\sigma : F(\sigma) \rightarrow H$ dados por $\varphi_\sigma(\sigma^m) = [\varphi(\sigma)]^m$ para $m \in \mathbb{Z}$, esto pues $\varphi(\sigma)$ la imagen del generador de $F(\sigma)$ y eso es lo que define por completo un homomorfismo que tiene como dominio un grupo cíclico. Ahora para la familia indizada de grupos $(F(\sigma))_{\sigma \in S}$ existe una única familia de homomorfismos $(\varphi_\sigma)_{\sigma \in S}$ con sus respectivas inclusiones, por lo que, por 2.36 existe un único homomorfismo $\Phi : F(S) \rightarrow H$ tal que, para cada $\sigma \in S$ el siguiente diagrama conmuta

$$\begin{array}{ccc} F(S) & & \\ \uparrow \iota_\sigma & \searrow \Phi & \\ F(\sigma) & \xrightarrow{\varphi_\sigma} & H \end{array}$$

Como se verifica $\Phi \circ \iota_\sigma = \varphi_\sigma$ y si $j : S \rightarrow F(\sigma)$ es tal que $j(\sigma) = \sigma$ donde j_σ es la inclusión $\{\sigma\} \hookrightarrow F(\sigma)$. Usando esta notación tendremos $\iota = \iota_\sigma \circ j$ y $\varphi = \varphi_\sigma \circ j$,

por tanto

$$\Phi \circ \iota_\sigma \circ j = \varphi_\sigma \circ j,$$

$$\Phi \circ \iota = \varphi,$$

la cual es la condición buscada. La unicidad de Φ también se deriva de la unicidad en la propiedad característica del producto libre, aunque también puede verificarse de manera directa revisando las imágenes de los generadores de $F(S)$ bajo Φ . ■

En general, la definición de grupo libre que tomaremos será la siguiente:

2.42 Definición. Un grupo arbitrario G es llamado *grupo libre* si existe un conjunto $S \subseteq G$ tal que el homomorfismo $\Phi : F(S) \rightarrow G$ inducido por la inclusión $S \hookrightarrow G$, es un isomorfismo.

Un detalle a tener en cuenta es que un producto libre no necesariamente será un grupo libre, por ejemplo el grupo $\mathbb{Z}/2\mathbb{Z} * \mathbb{Z}/2\mathbb{Z}$ no es un grupo libre, pues en los grupos libres únicamente las identidades tienen orden finito y en este caso hay elementos como la clase del 1 módulo 2 identificada como un elemento del producto directo, que tienen orden 2.

Y para concluir, la definición principal de esta sección:

2.43 Definición. Si S es un conjunto arbitrario y R es un conjunto de elementos en el grupo libre $F(S)$. Llamaremos *presentación de grupos* al par (S, R) que define el grupo

$$\langle S | R \rangle = \frac{F(S)}{\text{Ncl}_{F(S)}(R)},$$

donde $\text{Ncl}_{F(S)}(R)$ denota el conjunto normal más pequeño (respecto a la relación de orden de contención) que contiene a R . A los elementos de S y R se les conoce como *generadores* y *relaciones*, respectivamente.

Una caracterización de $\text{Ncl}_{F(S)}(R)$ dada explícitamente por sus elementos y no tan inmediata es la siguiente:

$$\text{Ncl}_{F(S)}(R) = \{f_1^{-1}r_1^{\varepsilon_1}f_1 \cdots f_n^{-1}r_n^{\varepsilon_n}f_n : n \geq 0, \varepsilon_i = \pm 1, r_i \in R, f_i \in F(S)\}.$$

En particular, esta definición viene motivada por expresar un grupo en función de su conjunto de generadores y las relaciones entre estos generadores, las cuales son,

de cierto modo, una especie de reglas que verifican los generadores del grupo. Ciertamente, de ahí viene el nombre de $F(S)$ pues puede pensarse como un grupo “libre de restricciones” entre sus elementos. A partir de esta idea, podemos pensar que si existe un homomorfismo sobreyectivo $\psi : F(S) \rightarrow G$ a algún grupo G arbitrario (es natural pensar esto pues en el caso en el que $S \subseteq G$ genera a G no es muy difícil construirlo) si su kernel fuese el conjunto más pequeño que contiene a un conjunto R , por el primer teorema de isomorfismos tendremos que $\langle S|R \rangle$ es isomorfo a G , mediante el isomorfismo inducido $\bar{\psi}$. De esta manera, forzamos a R las relaciones que pueden ser del tipo AB^{-1} a ser neutros y por lo tanto a tener la igualdad $A = B$. A veces conviene directamente referirnos a $\langle S|R \rangle$ como una presentación de grupo en lugar de al par (S, R) .

2.44 Definición. Sea G un grupo arbitrario. Una *presentación para G* es una presentación de grupo (S, R) tal que $\langle S|R \rangle \cong G$.

2.45 Ejemplo. Los grupos cíclicos tienen dos posibles presentaciones, $\langle x|\emptyset \rangle$ si es infinito y $\langle x|x^n \rangle$ si es finito con n elementos. En caso de no tener relaciones podemos escribir $\langle x|\emptyset \rangle = \langle x \rangle = F(x)$, aunque hay que tener cuidado con no confundir $\langle x \rangle$ con algún otro grupo cíclico.

2.46 Ejemplo. El Grupo diédrico D_n tiene la presentación $\langle s, r | r^n, s^2, (sr)^2 \rangle$.

2.47 Ejemplo. Dos presentaciones $\langle S_1, R_1 \rangle$ y $\langle S_2, R_2 \rangle$ pueden ser equivalentes si $\langle S_1, R_1 \rangle \cong \langle S_2, R_2 \rangle$, por ejemplo $\langle a, b, c | abc \rangle$ es equivalente a $\langle a, b | \emptyset \rangle$, pues c puede ser obtenido mediante los otros generadores.

2.4. Teorema de Van-Kampen

Antes de pasar al teorema principal del capítulo veremos una motivación. Sea X un espacio topológico y sean $U, V \subseteq X$ dos conjuntos abiertos tales que su unión es X y cuya intersección es no vacía, si tomamos cualquier punto $p \in U \cap V$ como base, las cuatro inclusiones inducen cuatro homomorfismos de grupos.

$$\begin{array}{ccccc}
 & & U & & \pi_1(U, p) \\
 & \nearrow i & \searrow k & \nearrow i_* & \searrow k_* \\
 U \cap V & & X & \implies & \pi_1(U \cap V, p) \\
 & \searrow j & \nearrow l & \searrow j_* & \nearrow l_* \\
 & & V & & \pi_1(V, p)
 \end{array}$$

Ahora podemos agregar el producto libre dentro de nuestro diagrama y por la propiedad característica del producto libre, podemos asegurar la existencia y unicidad de un homomorfismo $\Phi : \pi_1(U, p) * \pi_1(V, p) \rightarrow \pi_1(X, p)$, tal que el siguiente diagrama conmuta:

$$\begin{array}{ccccc}
 & & \pi_1(U, p) & & \\
 & \nearrow i_* & \downarrow & \searrow k_* & \\
 \pi_1(U \cap V, p) & & \pi_1(U, p) * \pi_1(V, p) & \xrightarrow{\Phi} & \pi_1(X, p) \\
 & \searrow j_* & \uparrow & \nearrow l_* & \\
 & & \pi_1(V, p) & &
 \end{array}$$

Y ahora podríamos intentar mostrar que Φ es sobreyectivo y caracterizar $\text{Ker } \Phi$ para buscar un isomorfismo de grupos mediante el primer teorema de isomorfismos. Sorprendentemente, no solo es posible mostrar la sobreyectividad de Φ , si no también describir su kernel de manera precisa y es uno de los métodos de cálculo más efectivos para grupos fundamentales.

2.48 Teorema (Seifert-Van Kampen). *Si X es un espacio topológico, supongase que $U, V \subseteq X$ son conjuntos abiertos tales que su unión es X , con U, V y $U \cap V$ conexos por caminos. Si $p \in U \cap V$, y Φ es tal que el siguiente diagrama conmuta:*

$$\begin{array}{ccccc}
 & & \pi_1(U, p) & & \\
 & \nearrow i_* & \downarrow \iota_{\pi_1(U, p)} & \searrow k_* & \\
 \pi_1(U \cap V, p) & & \pi_1(U, p) * \pi_1(V, p) & \xrightarrow{\Phi} & \pi_1(X, p) \\
 & \searrow j_* & \uparrow \iota_{\pi_1(V, p)} & \nearrow l_* & \\
 & & \pi_1(V, p) & &
 \end{array}$$

Con $i = \iota_{U \cap V}, j = j_{U \cap V}, k = \iota_U, l = \iota_V, \iota_{\pi_1(U, p)}$ y $\iota_{\pi_1(V, p)}$ mapeos inclusión. Entonces Φ es un homomorfismo sobreyectivo con:

$$\text{Ker } \Phi = \text{Ncl}_{\pi_1(U, p) * \pi_1(V, p)} \left(\{ (i_* \gamma)(j_* \gamma)^{-1} : \gamma \in \pi_1(U \cap V, p) \} \right).$$

Y por lo tanto:

$$\pi_1(X, p) \cong \frac{\pi_1(U, p) * \pi_1(V, p)}{\text{Ncl}_{\pi_1(U, p) * \pi_1(V, p)} \left(\{ (i_* \gamma)(j_* \gamma)^{-1} : \gamma \in \pi_1(U \cap V, p) \} \right)}.$$

La prueba de este teorema es bastante técnica y requiere definir un concepto importante dentro de la topología de variedades llamado *Número de Lebesgue*. Sin

embargo, es posible seguir casi toda la demostración con las herramientas desarrolladas hasta ahora y se invita al lector a consultarla en [19, pp. 268-273]. Para fines prácticos, daremos una prueba al final de la sección para un caso particular del teorema que nos ayudará a trabajar con espacios en los que estamos interesados, sin embargo, el teorema se enuncia por su valor dentro de la materia. Por ahora haremos algunos comentarios relacionados a los grupos que aparecen dentro del teorema, pues en general el cociente definido tiene un nombre:

2.49 Definición. Supongamos H, G_1 y G_2 grupos, con $f_1 : H \rightarrow G_1$ y $f_2 : H \rightarrow G_2$ homomorfismos. El *producto libre amalgamado de G_1 y G_2 sobre H* , denotado por $G_1 *_H G_2$ se define como

$$G_1 *_H G_2 = \frac{G_1 * G_2}{\text{Ncl}_{G_1 * G_2}(\{f_1(h)f_2(h)^{-1} : h \in H\})}.$$

Y un resultado muy útil para el cálculo de productos libres amalgamados es el siguiente teorema:

2.50 Teorema (Presentación para un producto libre amalgamado). *Si $f_1 : H \rightarrow G_1$ y $f_2 : H \rightarrow G_2$ son homomorfismos de grupos. Suponga que G_1, G_2 , y H tienen las siguientes presentaciones:*

$$\begin{aligned} G_1 &\cong \langle \alpha_1, \dots, \alpha_m | \rho_1, \dots, \rho_r \rangle; \\ G_2 &\cong \langle \beta_1, \dots, \beta_n | \sigma_1, \dots, \sigma_s \rangle; \\ H &\cong \langle \gamma_1, \dots, \gamma_p | \tau_1, \dots, \tau_t \rangle. \end{aligned}$$

Entonces el producto libre amalgamado tiene la siguiente presentación:

$$G_1 *_H G_2 \cong \langle \lambda_1, \dots, \lambda_{m+n} | \rho_1, \dots, \rho_r, \sigma_1, \dots, \sigma_s, u_1 v_1^{-1}, \dots, u_p v_p^{-1} \rangle,$$

donde $\lambda_i = \alpha_i$ para $1 \leq i \leq m$, $\lambda_i = \beta_i$ para $m+1 \leq i \leq m+n$, u_j es una expresión para $f_1(\gamma_j)$ en terminos de los generadores de la presentación de G_1 y v_j es una expresión para $f_2(\gamma_j)$ en terminos de los generadores de la presentación de G_2 .

Demostración. La prueba se divide en dos pasos, el primero es mostrar que si S_1, S_2 son conjuntos disjuntos y R_1, R_2 son subconjuntos de $F(S_1)$ y $F(S_2)$, respectivamente, entonces $\langle S_1 \cup S_2 | R_1 \cup R_2 \rangle \cong \langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle$ y el segundo paso es probar que si R y R' son dos subconjuntos de $F(S)$ para algún conjunto S , entonces $\langle S | R \cup R' \rangle \cong \langle S | R \rangle / \text{Ncl}_{\langle S | R \rangle}(\pi(R'))$, donde $\pi : F(S) \rightarrow \langle S | R \rangle$ es la proyección canónica. Una vez con estos hechos, el teorema es una consecuencia directa.

Primer paso: De 2.37, 2.41 y $S_1 \cap S_2 = \emptyset$, se sigue $F(S_1) * F(S_2) \cong F(S_1 \cup S_2)$ mediante un isomorfismo φ en el cual podemos abusar un poco de la notación y escribir como¹ $\varphi(x) = x$. Si $\psi_1 : F(S_1) \rightarrow \langle S_1 | R_1 \rangle$ y $\psi_2 : F(S_2) \rightarrow \langle S_2 | R_2 \rangle$ son tales que $\text{Ker } \psi_i = \text{Ncl}_{F(S_i)}(R_i)$ para $i = 1, 2$ (las proyecciones canónicas) existe un único homomorfismo Φ tal que

$$\begin{array}{ccc} F(S_1) & \xrightarrow{\psi_1} & \langle S_1 | R_1 \rangle \\ \downarrow \iota_1 & & \downarrow \iota'_1 \\ F(S_1) * F(S_2) & \xrightarrow{\phi} & \langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle \\ \uparrow \iota_2 & & \uparrow \iota'_2 \\ F(S_2) & \xrightarrow{\psi_2} & \langle S_2 | R_2 \rangle, \end{array}$$

pues esta condición es equivalente a 2.36 ya que se puede reescribir como:

$$\begin{array}{ccc} F(S_1) * F(S_2) & & F(S_1) * F(S_2) \\ \uparrow \iota_1 & \searrow \phi & \uparrow \iota_2 \\ F(S_1) & \xrightarrow{\iota'_1 \circ \psi_1} & \langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle. \end{array} \quad \begin{array}{ccc} F(S_1) * F(S_2) & & F(S_1) * F(S_2) \\ \uparrow \iota_2 & \searrow \phi & \uparrow \iota_1 \\ F(S_2) & \xrightarrow{\iota'_2 \circ \psi_2} & \langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle. \end{array}$$

Notemos que ϕ es sobreyectivo dado que ψ_1, ψ_2 son sobreyectivos. Por tanto si $\varphi : F(S_1 \cup S_2) \rightarrow F(S_1) * F(S_2)$ es el isomorfismo construido al inicio, y definimos el homomorfismo Φ tal que:

$$\begin{array}{ccccc} & & \Phi := \phi \circ \varphi & & \\ & \searrow & & \searrow & \\ F(S_1 \cup S_2) & \xrightarrow{\varphi} & F(S_1) * F(S_2) & \xrightarrow{\phi} & \langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle. \end{array}$$

Veremos que $\text{Ker } \Phi = \text{Ncl}_{F(S_1 \cup S_2)}(R_1 \cup R_2)$, omitiremos algunas inclusiones, sin embargo, es un detalle a tomar en cuenta. Es claro que $R_1 \cup R_2 \subseteq \text{Ker } \Phi$ pues si $r \in R_1 \cup R_2$ tomaremos $r \in R_1$ sin perdida de generalidad y escribiendo la identificación $\iota_1(r) = \varphi(r) \in R_1 \subseteq F(S_1)$ se tiene, pues $r \in \text{Ker } \psi_1$ lo siguiente:

$$\Phi(r) = \phi \circ \varphi(r) = \phi \circ \iota_1(r) = \iota'_1 \circ \psi_1(r) = 1_{\langle S_1 | R_1 \rangle} = 1_{\langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle}.$$

Como r arbitrario se tiene $R_1 \cup R_2 \subseteq \text{Ker } \Phi$ y por la minimalidad del grupo normal

¹Esto es similar a lo que ocurre con productos directos, si G_1, G_2 y G_3 son grupos arbitrarios es posible probar que $(G_1 \times G_2) \times G_3 \cong G_1 \times (G_2 \times G_3)$, pero usualmente se escriben como un mismo grupo $G_1 \times G_2 \times G_3$.

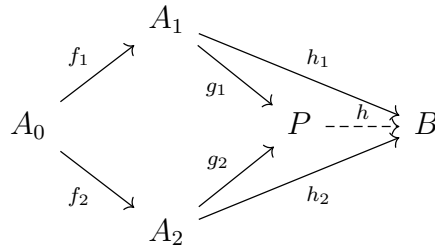
que estamos trabajando $\text{Ncl}_{F(S_1 \cup S_2)}(R_1 \cup R_2) \subseteq \text{Ker } \Phi$. Para la otra parte supongamos que $\Phi(s) = 1_{\langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle}$, escribimos $s = r_1^{m_1} \mathbf{r}_1^{n_1} \dots r_\ell^{m_\ell} \mathbf{r}_\ell^{n_\ell}$ donde $r_j \in S_1$ y $\mathbf{r}_j \in S_2$ para $j = 1, \dots, \ell$, entonces

$$\psi_1(r_1^{m_1})\psi_2(\mathbf{r}_1^{n_1}) \dots \psi_1(r_\ell^{m_\ell})\psi_2(\mathbf{r}_\ell^{n_\ell}) = \Phi(s) = 1_{\langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle},$$

por lo que esto implica $r_j^{m_j} \in \text{Ker } \psi_1 = \text{Ncl}_{F(S_1)}(R_1)$ y $\mathbf{r}_j^{n_j} \in \text{Ker } \psi_2 = \text{Ncl}_{F(S_2)}(R_2)$, luego cada $r_j^{m_j}$ es un elemento de R_1 o un producto de elementos en R_1 con elementos de $F(S_1)$ (mutatis mutandis para los $\mathbf{r}_j^{n_j}$) los cuales al multiplicarlos todos se obtiene que $s \in \text{Ncl}_{F(S_1 \cup S_2)}(R_1 \cup R_2)$, lo cual muestra $\text{Ker } \Phi \subseteq \text{Ncl}_{F(S_1 \cup S_2)}(R_1 \cup R_2)$. Finalmente, por el primer teorema de isomorfismo $\Phi : F(S_1 \cup S_2) \rightarrow \langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle$ desciende a un isomorfismo $\tilde{\Phi} : \langle S_1 \cup S_2 | R_1 \cup R_2 \rangle \rightarrow \langle S_1 | R_1 \rangle * \langle S_2 | R_2 \rangle$.

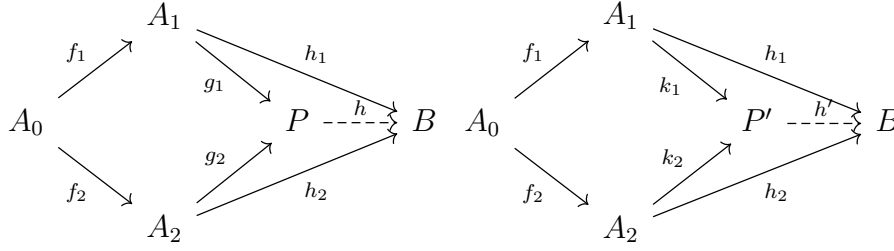
Segundo paso: Consideremos el homomorfismo conocido como la proyección canónica $\phi : \langle S | R \rangle \rightarrow \langle S | R \rangle / \text{Ncl}_{\langle S | R \rangle}(\pi(R'))$ y si $\phi \circ \pi : F(S) \rightarrow \langle S | R \rangle / \text{Ncl}_{\langle S | R \rangle}(\pi(R'))$, al tener la composición de dos homomorfismos sobreyectivos tendremos otro homomorfismo sobreyectivo el cual puede verificarse que tiene kernel $\text{Ncl}_{F(S)}(R \cup R')$ y el resultado se sigue de nueva cuenta por el primer teorema de isomorfismos. ■

Ahora, para quien interese, daremos una mirada al producto libre amalgamado desde un punto de vista más categórico. Si A_0, A_1 y A_2 son objetos en una categoría $\mathbf{C}$ y sean $f_i \in \text{Hom}_{\mathbf{C}}(A_0, A_i)$ para $i = 1, 2$. Un *Pushout de la pareja* (f_1, f_2) es un objeto $P \in \text{Ob}(\mathbf{C})$ junto con un par de morfismos $g_i \in \text{Hom}_{\mathbf{C}}(A_i, P)$ tal que $g_1 \circ f_1 = g_2 \circ f_2$ que satisface la siguiente propiedad característica: dado un objeto $B \in \text{Ob}(\mathbf{C})$ y morfismos $h_i \in \text{Hom}_{\mathbf{C}}(A_i, B)$ tales que $h_1 \circ f_1 = h_2 \circ f_2$, existe un único morfismo $h \in \text{Hom}_{\mathbf{C}}(P, B)$ tal que $h \circ g_i = h_i$ para $i = 1, 2$:

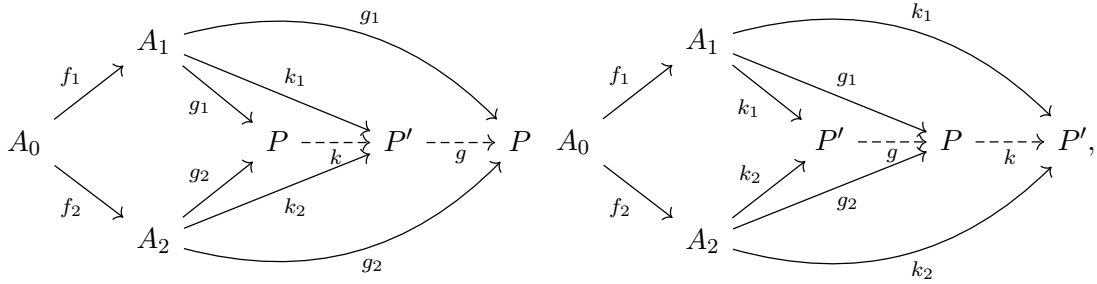


Veamos que si el Pushout para una pareja de morfismos existe, entonces este es único salvo isomorfismos en $\mathbf{C}$. Para ello supongamos que (P, g_1, g_2) y (P', k_1, k_2) son

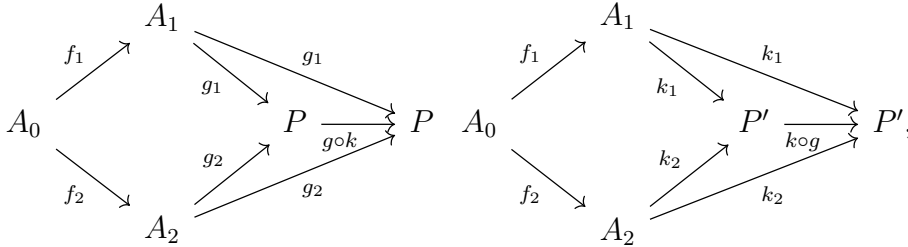
Pushouts, entonces existen morfismos únicos h, h' tales que los siguientes diagramas



conmutan para cualquier $B \in \text{Ob}(\mathcal{C})$ y morfismos $h_i \in \text{Hom}_{\mathcal{C}}(A_i, B)$. Por tanto existen g, k morfismos tales que los siguientes diagramas conmutan:



para dar lugar a los siguientes diagramas:



que al ser conmutativos incluso cuando se reemplaza Id_P y $\text{Id}_{P'}$ por $g \circ k$ y $k \circ g$, respectivamente, se concluye por la unicidad en la definición de pushout, que $\text{Id}_P = g \circ k$ y $\text{Id}_{P'} = k \circ g$. De esta manera, el producto libre amalgamado puede ser descrito en estos términos.

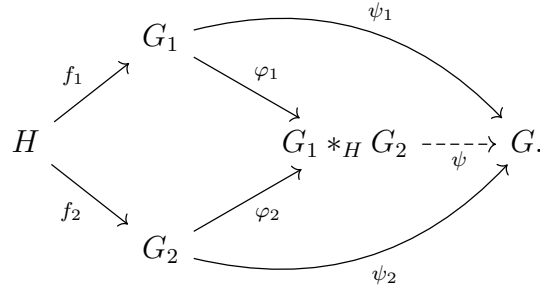
2.51 Proposición. *El producto libre amalgamado es el Pushout de dos homomorfismos con un mismo dominio, es decir, un Pushout en la categoría Grp.*

Demostración. Dado un par $f_1: H \rightarrow G_1$ y $f_2: H \rightarrow G_2$ de homomorfismos, si tomamos $\iota_{G_1}: G_1 \rightarrow G_1 * G_2$ y $\iota_{G_2}: G_2 \rightarrow G_1 * G_2$ los homomorfismos dados por las inclusiones y $\pi: G_1 * G_2 \rightarrow G_1 *_H G_2$ es la proyección canónica al cociente con $\text{Ker } \pi = \text{Ncl}_{G_1 * G_2}(\{f_1(h)f_2(h)^{-1} : h \in H\}) = N$. Si $\varphi_1 := \pi \circ \iota_{G_1}$ y $\varphi_2 := \pi \circ \iota_{G_2}$

notemos que, si $h \in H$, entonces $f_1(h) \hookrightarrow \iota_{G_1}(f_1(h))$ y $f_2(h) \hookrightarrow \iota_{G_2}(f_2(h)^{-1})$ y se tiene:

$$\begin{aligned}\pi(\iota_{G_1}(f_1(h))\iota_{G_2}(f_2(h)^{-1})) &= 1_{G_1 *_H G_2}, \\ \pi(\iota_{G_1}(f_1(h)))\pi(\iota_{G_2}(f_2(h)))^{-1} &= 1_{G_1 *_H G_2}, \\ \pi \circ \iota_{G_1}(f_1(h)) &= \pi \circ \iota_{G_2}(f_2(h)), \\ \varphi_1 \circ f_1(h) &= \varphi_2 \circ f_2(h).\end{aligned}$$

Por lo tanto $\varphi_1 \circ f_1 = \varphi_2 \circ f_2$. Si suponemos (ψ_1, ψ_2, G) otros homomorfismos y un grupo G , tales que $\psi_1 : G_1 \rightarrow G$ y $\psi_2 : G_2 \rightarrow G$ con $\psi_1 \circ f_1 = \psi_2 \circ f_2$, mostraremos que existe un único $\psi : G_1 *_H G_2 \rightarrow G$ tal que el siguiente diagrama conmuta:



Para ello usaremos la Proposición 2.36, conocida como propiedad universal del producto libre, notemos que esta propiedad nos garantiza que un único $p : G_1 * G_2 \rightarrow G$ tal que los siguientes diagramas conmutan:

$$\begin{array}{ccc} G_1 * G_2 & & G_1 * G_2 \\ \iota_{G_1} \uparrow & \searrow p & \uparrow \iota_{G_2} \\ G_1 & \xrightarrow{\psi_1} & G \end{array} \quad \begin{array}{ccc} G_1 * G_2 & & G_1 * G_2 \\ \iota_{G_2} \uparrow & \searrow p & \uparrow \iota_{G_1} \\ G_2 & \xrightarrow{\psi_2} & G \end{array}$$

el cual verifica $N \subseteq \text{Ker } p$ pues si $h \in H$ se cumple $\psi_1(f_1(h))\psi_2(f_2(h)^{-1}) = 1_G$ por tanto $p(\iota_{G_1}(f_1(h))\iota_{G_2}(f_2(h)^{-1})) = 1_G$ y de esto $p(\iota_{G_1}(f_1(h))\iota_{G_2}(f_2(h)^{-1})) = 1_G$, así que todos los elementos de la forma $\iota_{G_1}(f_1(h))\iota_{G_2}(f_2(h)^{-1})$ están en el kernel, por la minimalidad de N sigue la contención. La relación $N \subseteq \text{Ker } p$ nos asegura que la siguiente función está bien definida y es un homomorfismo que cumple nuestros requisitos:

$$\begin{aligned}\psi : G_1 *_H G_2 &\rightarrow G, \\ gN &\mapsto p(g).\end{aligned}$$

Además es único pues depende únicamente de los valores que toma en G_1, G_2 (con sus respectivas inclusiones) y allí debe coincidir con ψ_1, ψ_2 , respectivamente, por lo que está condicionado por esta restricción. ■

2.52 Corolario. *El producto libre amalgamado es único salvo isomorfismos.*

Antes de pasar a la prueba del caso particular de Van-Kampen consideraremos el siguiente lema. Si el lector no está familiarizado con la terminología relativa a espacios de recubrimiento se recomienda consultar los capítulos 11 y 12 en [19].

2.53 Lema. *Sean M_1 y M_2 conjuntos abiertos y conexos cubriendo una variedad topológica M con $M_1 \cap M_2$ conexo y no vacío. Si $\phi_1 : N_1 \rightarrow M_1$ y $\phi_2 : N_2 \rightarrow M_2$ son aplicaciones de recubrimiento que se vuelven isomorfas al restringir la base a $M_1 \cap M_2$, entonces existe un recubrimiento $\phi : N \rightarrow M$ que se vuelve isomorfo a ϕ_1 restringiendo la base a M_1 , y se vuelve isomorfo a ϕ_2 restringiendo la base a M_2 .*

Demostración. Como ϕ_1, ϕ_2 son isomorfos al restringir la base a $M_1 \cap M_2$, existe $\varphi : \phi_1^{-1}(M_1 \cap M_2) \rightarrow N_2$ tal que $\phi_2 \circ \varphi = \phi_1$ (restringiendo las aplicaciones de recubrimiento), por lo que se define $N = N_2 \cup_{\varphi} N_1$ y la función $\phi : N \rightarrow M$ como:

$$\phi([(x, i)]) := \begin{cases} \phi_1(x) & \text{si } x \in N_1, \\ \phi_2(x) & \text{si } x \in N_2. \end{cases}$$

Notemos que está bien definida, pues si tuviéramos $[(x, i)] = [(y, j)]$ hay que verificar 4 casos.

- i) Si $x \in N_1 \setminus \phi_1^{-1}(M_1 \cap M_2)$, entonces $[(x, 1)] = \{(x, 1)\}$, por tanto $x \in N_1$ y sin ambigüedad $\phi([(x, 1)]) = \phi_1(x)$.
- ii) Como segundo caso $x \in \phi_1^{-1}(M_1 \cap M_2)$ entonces $[(x, 1)] = \varphi^{-1}(\{\varphi(x)\}) \cup \{\varphi(x)\}$ y puede pasar que $[(x, 1)] = [(y, 1)]$ o $[(x, 1)] = [(y, 2)]$. Todo $y \in N_1$ que verifique $[(y, 1)] = [(x, 1)]$ es tal que $\varphi(y) = \varphi(x)$, de lo que tendremos

$$\phi_1(y) = \phi_2 \circ \varphi(y) = \phi_2 \circ \varphi(x) = \phi_1(x).$$

La otra opción del segundo caso (donde $[(x, 1)] = [(y, 2)]$) es que $y = \varphi(x)$ y aquí, por la condición de isomorfismo, al hacer la composición se sigue que $\phi_2(y) = \phi_1(x)$, por lo que siempre se tiene $\phi([(y, j)]) = \phi([(x, 1)])$.

- iii) En el tercer caso si $x \in N_2 \setminus \varphi(\phi_1^{-1}(M_1 \cap M_2))$, entonces $[(x, 2)] = \{(x, 2)\}$ y no hay ambigüedad.
- iv) Para el último caso, cuando $x \in \varphi(\phi_1^{-1}(M_1 \cap M_2))$ se tiene $x = \varphi(a)$ para algún $a \in \phi_1^{-1}(M_1 \cap M_2)$ por lo que $[(x, 2)] = [(a, 1)]$ y la verificación pasa a ser el caso ii).

La sobreyectividad de ϕ viene de que ϕ_1, ϕ_2 son sobreyectivas y $M = M_1 \cup M_2$. Para mostrar continuidad basta probar que la composición $\phi \circ q : N_2 \sqcup N_1 \rightarrow M$ con $q : N_2 \sqcup N_1 \rightarrow N_2 \cup_\varphi N_1$, la aplicación cociente, es continua. Para esto es suficiente mostrar que $\phi \circ q \circ \iota_{N_i} = \phi \circ q|_{N_i} : N_i \rightarrow M$ son continuas para $i = 1, 2$; sin embargo, esto es consecuencia de que:

$$\phi \circ q|_{N_i} = \phi \circ q \circ \iota_{N_i} = \iota_{M_i} \circ \phi_i,$$

y el lado derecho de la igualdad es una función continua para $i = 1, 2$; por lo tanto $\phi \circ q$ es continua. Como M_1, M_2 son subconjuntos abiertos de una variedad M , estos conjuntos son variedades de la misma dimensión que M y al ser N_1, N_2 sus espacios de recubrimiento respectivamente, también son variedades y además son conexas. Sumado a que las variedades son localmente conexas por caminos, los espacios N_1, N_2 serán conexos por caminos. Por todo esto $N_2 \cup_\varphi N_1$ también es un espacio conexo por caminos (por ende conexo y localmente conexo por caminos) cuando $\phi_1^{-1}(M_1 \cap M_2) \neq \emptyset$. Por último, para ser una aplicación de recubrimiento, falta mostrar que M tiene una base uniformemente cubierta por ϕ , que se deriva de las propiedades de ϕ_1 y ϕ_2 como aplicaciones de recubrimiento. ■

La siguiente versión del teorema de Seifert-Van Kampen es más débil pues es válida únicamente para variedades, sin embargo, en la práctica la mayoría de espacios interesantes son variedades o espacios con el mismo tipo de homotopía que una variedad. Fulton en [12, p. 199] atribuye algunas de estas ideas a Grothendieck, sin embargo, el enfoque que usaremos será el de Terence Tao en [36], quien en 2012 dio algunos detalles e ideas extra en esta prueba.

2.54 Teorema (Seifert-Van Kampen, versión para variedades). *Sean M_1, M_2 conjuntos abiertos y conexos cubriendo una variedad topológica M con $M_1 \cap M_2$ también conexo, si p es un punto de $M_1 \cap M_2$ entonces*

$$\pi_1(M_1 \cup M_2, p) \cong \pi_1(M_1, p) *_{\pi_1(M_1 \cap M_2, p)} \pi_1(M_2, p).$$

Demostración. Sea G un grupo. Supongamos que existen homomorfismos de grupos $f_1 : \pi_1(M_1, p) \rightarrow G$, $f_2 : \pi_1(M_2, p) \rightarrow G$, tales que el siguiente diagrama conmuta para ι_1, ι_2 homomorfismos canónicos:

$$\begin{array}{ccccc}
 & & \pi_1(M_1, p) & & \\
 & \nearrow \iota_1 & & \searrow f_1 & \\
 \pi_1(M_1 \cap M_2, p) & & & & G \\
 & \searrow \iota_2 & & \nearrow f_2 & \\
 & & \pi_1(M_2, p) & &
 \end{array}$$

Mostraremos que existe un único homomorfismo $f : \pi_1(M, p) \rightarrow G$ tal que los f_i verifican $f \circ j_i = f_i$ donde $j_i : \pi_1(M_i, p) \rightarrow \pi_1(M_1 \cup M_2, p)$ son los homomorfismos canónicos para $i = 1, 2$. De esta manera, obtendremos que $\pi_1(M, p) = \pi_1(M_1 \cup M_2, p)$ es un pushout y por el Corolario 2.52 nuestra prueba estará completa.

Para la existencia. Tomaremos $\tilde{\phi}_1 : \tilde{M}_1 \rightarrow M_1$ el recubrimiento universal para M_1 , por tanto $\tilde{M}_1$ es simplemente conexo, y si tomamos un punto base $x_1 \in \tilde{\phi}_1^{-1}(\{p\})$, entonces para cualquier otro punto y en el fibrado, existe un único elemento $g \in \pi_1(M_1, p)$ tal que $y = x_1 g$, donde $x_1 g$ es la acción por la derecha de g por monodromía en x_1 . Esto da la acción por la izquierda de $\pi_1(M_1, p)$ en $\tilde{M}_1$ por automorfismos de recubrimiento o “transformaciones deck” $\tau_h : \tilde{M}_1 \rightarrow \tilde{M}_1$ para cada $h \in \pi_1(M_1, p)$ dadas por

$$\tau_h(x_1 g) = x_1 h g.$$

Las fibras del recubrimiento universal $\tilde{\phi}_1$ son copias de $\pi_1(M_1, p)$. Con esto podremos formar un nuevo recubrimiento $\phi_1 : N_1 \rightarrow M_1$ cuyas fibras son copias de G (con la topología discreta), formando primero el producto cartesiano $G \times \tilde{M}_1$ (que también es un recubrimiento de M_1) y luego tomando el cociente con la relación de equivalencia $(g f_1(g_1), x) \sim (g, \tau_{g_1} x)$ para todo $g \in G$, $x \in \tilde{M}_1$, $g_1 \in \pi_1(M_1, p)$. Este es un espacio de recubrimiento (posiblemente desconexo) de M_1 , cuyas fibras sobre p pueden ser identificadas con G identificando $g \in G$ con la clase de equivalencia $[(g, x_1)]$. La acción por monodromía (por la derecha) de $\pi_1(M_1, p)$ en esta fibra es entonces identificada con la acción por la derecha de $\pi_1(M_1, p)$ en G inducida por f_1 . Más aún, el grupo G actúa en N_1 por la izquierda vía automorfismos de recubrimiento, con cada $g \in G$ mapeando $[(g', x)]$ a $[(g g', x)]$.

Similarmente, podemos formar un recubrimiento $\phi_2 : N_2 \rightarrow M_2$ de M_2 cuya acción por monodromía de $\pi_1(M_2, p)$ en la fibra $\phi_2^{-1}(\{p\})$ pueden ser identificadas con la acción por la derecha de $\pi_1(M_2, p)$ en G inducida por f_2 . Cuando ambos recubrimientos se restringen a $M_1 \cap M_2$, la acción por monodromía de $\pi_1(M_1 \cap M_2, p)$ en la fibra sobre p , serán isomorfas entre ellas y las restricciones también son isomorfas entre ellas. Por el Lema 2.53, podemos pegar estos dos espacios de recubrimiento y obtener una aplicación de recubrimiento $\phi : N \rightarrow M$ isomorfa a ϕ_1 o ϕ_2 dependiendo de si la base se restringe a M_1 o M_2 , de tal manera que los cuatro isomorfismos forman un diagrama conmutativo; en particular, la fibra $\phi^{-1}(\{p\})$ aún está definida por G . La acción por monodromía de $\pi_1(M, p)$ en $\phi^{-1}(\{p\})$ se restringe a la previamente descrita de $\pi_1(M_1, p)$ en G y de $\pi_1(M_2, p)$ en G . Además, como G también actúa por la izquierda por automorfismos de recubrimiento en N_1 y N_2 , de tal forma que pueden ser vistos como compatibles con la restricción del isomorfismo a $M_1 \cap M_2$, G continua actuando por automorfismos de recubrimiento en el recubrimiento pegado N . Como las acciones de recubrimiento conmutan con los automorfismos de recubrimiento, concluimos que la acción por la derecha de un elemento g de $\pi_1(M_1, p)$ en G esta dado por la multiplicación por la derecha de un elemento $f(g)$ en G . Y al expandir las cuentas tendremos que f es un homomorfismo con las propiedades que buscamos.

Finalmente para la unicidad. Basta mostrar que $\pi_1(M, p)$ es generado por las imágenes de $\pi_1(M_1, p)$ y $\pi_1(M_2, p)$. Supongamos que esto no es cierto, entonces las imágenes de $\pi_1(M_1, p)$ y $\pi_1(M_2, p)$ generan un subgrupo propio G de $\pi_1(M, p)$. Sea $\tilde{\phi} : \tilde{M} \rightarrow M$ el recubrimiento universal de M , tal que $\tilde{\phi}^{-1}(\{p\})$ puede ser identificado con $\pi_1(M, p)$ (luego de fijar un punto en la fibra con se mencionó antes). Si $\tilde{M}_1$ es la restricción de $\tilde{M}$ a M_1 . La acción por monodromía (por la derecha) de $\pi_1(M, p)$ en la fibra encima de p , es inducida por el homomorfismo canónico de $\pi_1(M_1, p)$ a $\pi_1(M, p)$. En particular, G es preservado por su acción, y por ello podemos encontrar un sub cubierta N_1 de $\tilde{M}_1$ cuyas fibras corresponden a G , de manera análoga con $\tilde{M}_2$ se construye un sub cubierta N_2 con la misma fibra. Por tanto por el Lema 2.53 se construye un sub cubierta N de $\tilde{M}$. Pero esto es absurdo pues $\tilde{M}$ es conexo, de esto, la unicidad es cierta. ■

Podemos notar que esta prueba aprovecha una propiedad anteriormente discutida, el pushout de grupos. Ciertamente la topología algebraica se encarga de

estudiar este tipo de relaciones, muchas veces aprovechando la generalidad de la teoría de categorías, siendo esta una de las motivaciones detrás de las ideas en esta prueba. Realmente el requisito de M_1, M_2 ser variedades solo se usa para que nuestros espacios sean localmente conexos por caminos y aprovechar las propiedades de los espacios de recubrimiento, pero estas condiciones pueden relajarse y cambiar las hipótesis para obtener una prueba más general.

3. INVARIANTES ALGEBRAICOS

Una técnica para diferenciar dos nudos es tratar de formular invariantes algebraicos que caractericen estos objetos. En este contexto consideraremos un invariante como una propiedad que se preserva bajo isotopías de ambiente u otras equivalencias topológicas que se han presentado ya. Un ejemplo de un invariante (no algebraico) es la tricolorabilidad de un nudo¹. Continuando directamente con las herramientas desarrolladas en el capítulo pasado y regresando a la teoría de nudos, presentaremos el grupo de un nudo y revisaremos una prueba rigurosa de la existencia de nudos no triviales. Posteriormente, revisaremos otros invariantes importantes dentro de la teoría de nudos como el corchete de Kauffman y la homología de Khovanov.

3.1. Grupo de un nudo

Habiendo desarrollado nuestras herramientas algebraicas para el cálculo de grupos fundamentales, podríamos intentar usarlas para clasificar nudos, sin embargo, hay que ser bastante cuidadosos pues si $K \subset \mathbb{R}^3$ es un nudo, podríamos considerar $\pi_1(K, p)$ para algún $p \in K$ y notar que esta construcción no es muy interesante, pues K se definió como un conjunto homeomorfo a $\mathbb{S}^1$, de lo cual obtenemos por la Proposición 2.18 que $\pi_1(K, p) \cong \mathbb{Z}$ para cualquier nudo. Motivados por esto y recordando que la teoría de nudos se enfoca más en la topología de una variedad de dimensión 3, la siguiente definición resulta más útil.

3.1 Definición. Dado un nudo $K \subset \mathbb{R}^3$ y un punto $p \in \mathbb{R}^3 \setminus K$ al grupo fundamental $\pi_1(\mathbb{R}^3 \setminus K, p)$ se le conoce como el *grupo del nudo K con base p* o simplemente el *grupo de K* .

De manera similar, podemos sustituir $\mathbb{R}^3$ por $\mathbb{S}^3$ en la definición anterior y a $\pi_1(\mathbb{S}^3 \setminus K, p)$ también lo llamaremos grupo del nudo K , siempre aclararemos a cual de las dos opciones nos estamos refiriendo. Aprovechando que tanto $\mathbb{R}^3 \setminus K$ como

¹Este invariante también puede ser usado para verificar que el nudo trébol no es equivalente al nudo trivial, para más detalles ver [18, pp. 21-24].

$\mathbb{S}^3 \setminus K$ son conexos por caminos podemos omitir el punto base y simplemente hablar de $\pi_1(\mathbb{R}^3 \setminus K)$ y $\pi_1(\mathbb{S}^3 \setminus K)$. De la definición 3.1 se deriva directamente el siguiente resultado:

3.2 Teorema. *Dos nudos equivalentes tienen grupos isomorfos, más precisamente, si K_1, K_2 son dos nudos equivalentes, entonces $\pi_1(\mathbb{R}^3 \setminus K_1) \cong \pi_1(\mathbb{R}^3 \setminus K_2)$.*

Demostración. Supongamos que K_1, K_2 son dos nudos equivalentes entonces existe un homomorfismo $f : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ tal que $f(K_1) = K_2$ eso quiere decir que si $g = f|_{\mathbb{R}^3 \setminus K_1} : \mathbb{R}^3 \setminus K_1 \rightarrow \mathbb{R}^3 \setminus K_2$, la función g es un homomorfismo y se tiene que para todo $p \in \mathbb{R}^3 \setminus K_1$, por la Proposición 2.18, $\pi_1(\mathbb{R}^3 \setminus K_1, p) \cong \pi_1(\mathbb{R}^3 \setminus K_2, g(p))$. ■

Tomando en cuenta la definición de equivalencia entre nudos en $\mathbb{S}^3$, el resultado también es cierto para nudos dentro de este espacio. También es posible definir el grupo para un enlace.

3.3 Definición. Dado un enlace $L \subset \mathbb{R}^3$ y un punto $p \in \mathbb{R}^3 \setminus L$ al grupo fundamental $\pi_1(\mathbb{R}^3 \setminus L, p)$ se le conoce como el *grupo del enlace L con base p* o simplemente el *grupo de L* .

Y al igual que con los nudos, la definición es análoga para enlaces dentro $\mathbb{S}^3$. Con estas definiciones pasaremos a los ejemplos.

3.4 Ejemplo. El grupo del nudo trivial $\pi_1(\mathbb{R}^3 \setminus \mathbb{S}^1)$ no queda tan claro como calcularlo directamente, sin embargo se sabe que $\mathbb{R}^3 \setminus \mathbb{S}^1 \equiv \mathbb{S}^2 \vee \mathbb{S}^1$ así que por la Proposición 2.19 se tendrá que $\pi_1(\mathbb{R}^3 \setminus \mathbb{S}^1) \cong \pi_1(\mathbb{S}^2 \vee \mathbb{S}^1)$. Ahora usaremos el teorema de Van-Kampen para el calculo de este grupo fundamental. Primero consideremos $\mathbb{S}^2 \vee \mathbb{S}^1$ como subconjunto de $\mathbb{R}^3$ y con esto se consideran los abiertos U', V' que dan lugar a los abiertos U, V de la topología del subespacio en $\mathbb{S}^2 \vee \mathbb{S}^1$, tal y como en la Figura 3.1. Tomaremos $p \in U \cap V$ como el punto base de $\mathbb{S}^1$ y $\mathbb{S}^2$ en $\mathbb{S}^2 \vee \mathbb{S}^1$ y notemos que U tiene una parte contraíble y tiene el mismo tipo de homotopía que $\mathbb{S}^1$, V tiene el mismo tipo de homotopía que $\mathbb{S}^2$ y, finalmente, $U \cap V$ se puede contraer a p , así que tendremos $\pi_1(U, p) \cong \mathbb{Z}$, $\pi_1(V, p) = \{1_{\pi_1(V, p)}\}$, $\pi_1(U \cap V, p) = \{1_{\pi_1(U \cap V, p)}\}$ y por el teorema de Van-Kampen:

$$\pi_1(\mathbb{S}^2 \vee \mathbb{S}^1, p) \cong \pi_1(U, p) *_{\pi_1(U \cap V, p)} \pi_1(V, p) \cong \mathbb{Z}.$$

De lo que se concluye $\pi_1(\mathbb{R}^3 \setminus \mathbb{S}^1) \cong \mathbb{Z}$.

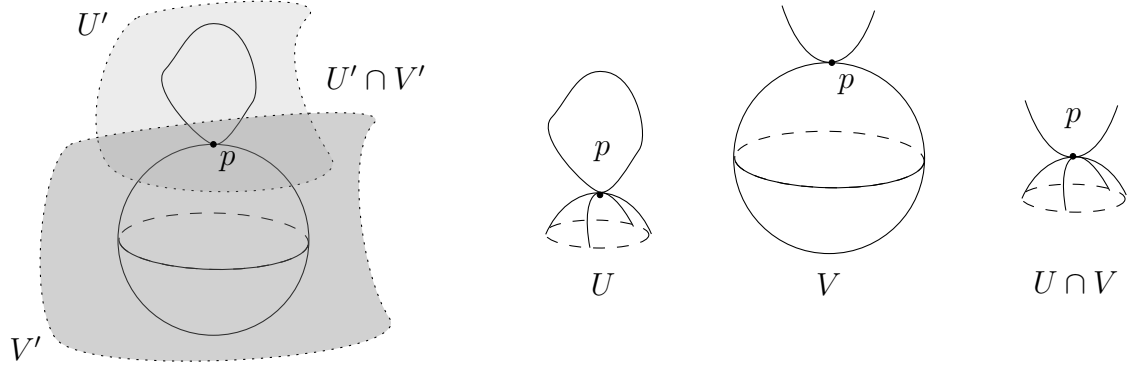


Figura 3.1. Aplicación de Van-Kampen en $\mathbb{S}^2 \vee \mathbb{S}^1$.

Para los detalles del porque $\mathbb{R}^3 \setminus \mathbb{S}^1 \equiv \mathbb{S}^2 \vee \mathbb{S}^1$, puede consultarse [13, p. 46]. Para describir grupos de nudos con más cruces o enlaces más complicados emplearemos la *presentación de Wirtinger* basándonos en el libro de Stillwell en [35, pp. 144-149]. Este método consiste en lo siguiente: al tomar el diagrama de un nudo, este puede ser dividido en una cantidad finita de componentes conexas A_i , donde i varía en el número de cruces. Si escogemos una orientación para el nudo y para los generadores a_i del grupo del nudo (tomando la convención de la mano derecha tal como un producto cruz) con el objetivo de simplificar la notación con índices y tener la componente A_i antes de la componente A_{i+1} marcando el cruce (por debajo) como la transición entre componentes, así como el generador $a_i = [\alpha_i]$ rodeando la componente A_i como en la Figura 3.2. Con esta notación podemos obtener relaciones entre los generadores, pues tal y como se muestra en la Figura 3.3 la composición de lazos asociada al producto de clases $a_i a_j^{-1} a_{i+1}^{-1} a_j$ puede contraerse al punto base.

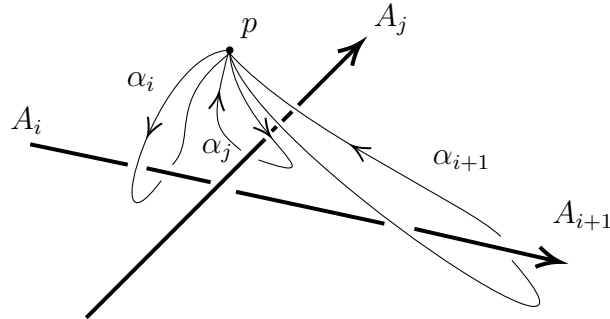


Figura 3.2. Generadores orientados en una vecindad de un cruce.

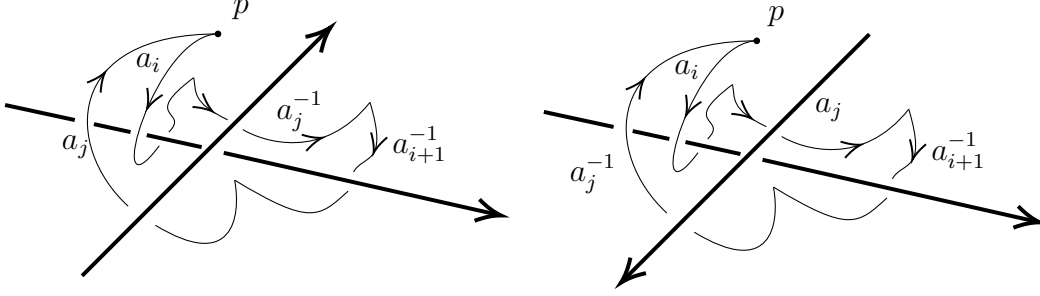


Figura 3.3. Lazos homotópicos al punto base en función de los generadores.

Y para la otra posibilidad del cruce, la relación entre los generadores sería $a_i a_j a_{i+1}^{-1} a_j^{-1}$, pues podemos tener tanto un cruce de la forma $\times$ como de la forma $\times$. Así, podemos enunciar la siguiente proposición.

3.5 Proposición (Presentación de Wirtinger). *Si K es un nudo cuyo diagrama tiene n cruces, entonces:*

$$\pi_1(\mathbb{R}^3 \setminus K) \cong \langle a_1, \dots, a_n | r_1, \dots, r_n \rangle,$$

donde cada r_j es de la forma $a_i a_j^{-1} a_\ell^{-1} a_j$ o $a_i a_j a_\ell^{-1} a_j^{-1}$.

Demostración. Esta proposición se deriva directamente del teorema de Van-Kampen. Usando la notación hasta ahora empleada, podemos asumir que cada una de las componentes conexas A_i del diagrama de nudo son segmentos en el plano $z = 1$ donde los puntos en la frontera (los extremos que van antes de los cruces) se unirán mediante una recta a su proyección en $z = 0$. Con esta descripción de los A_i podemos unir a A_i con A_{i+1} mediante un segmento contenido en $z = 0$ para obtener un nudo equivalente al que da origen al diagrama, así que sin pérdida de generalidad diremos que es el nudo K . Ahora removeremos de $\mathbb{R}^3$ una vecindad tubular $\mathcal{N}$ de K para la cual exista una deformación retráctil de $\mathbb{R}^3 \setminus K$ a $\mathbb{R}^3 \setminus \mathcal{N}$ como en la Figura 3.4.

Esta vecindad tubular puede ser vista como un “túnel rectangular” con cuadrados de lado ε para un épsilon suficientemente pequeño como secciones transversales. Consideramos $\mathcal{A} = \{z > 0\} \setminus \mathcal{N}$ y notemos que este conjunto tiene el mismo tipo de homotopía que $\bigvee_{i=1}^n \mathbb{S}^1$, por ende $\pi_1(\mathcal{A}) \cong \langle a_1, \dots, a_n | \emptyset \rangle$. Otra manera de ver esto es tomando un punto base y viendo que cada lazo puede o no estar debajo de los arcos rectangulares que se forman y de esta manera también es posible ver que el grupo fundamental de $\mathcal{A}$ es el propuesto. Por otro lado, $\mathcal{B} = \{\varepsilon/2 > z\} \setminus \mathcal{N}$ es un conjunto simplemente conexo. Finalmente, podemos notar que $\mathcal{A} \cap \mathcal{B} = \{\varepsilon/2 > z > 0\} \setminus \mathcal{N}$

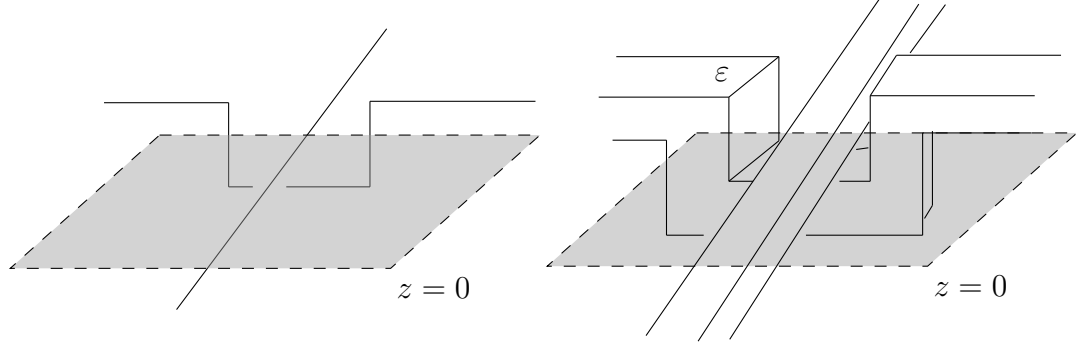


Figura 3.4. Túnel rectangular $\mathcal{N}$ usando un nudo como base.

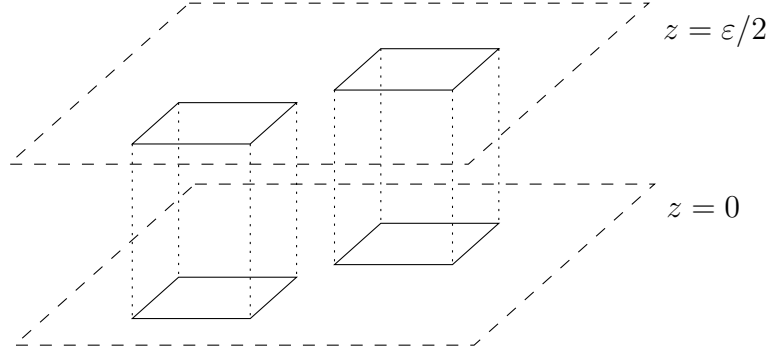


Figura 3.5. Intersección para el uso del teorema de Van-Kampen.

puede verse como una caja infinita con columnas n perforadas como en la Figura 3.5, por lo que de nueva cuenta $\pi_1(\mathcal{A} \cap \mathcal{B})$ es un grupo libre en n generadores y al identificarlos como elementos de $\pi_1(\mathcal{A})$ son elementos de la forma $a_i a_j^{-1} a_{i+1}^{-1} a_j$ o $a_i a_j a_{i+1}^{-1} a_j^{-1}$ dependiendo del tipo de cruce por el que pasan nuestros generadores. Con estos conjuntos se tendrá, por el teorema de Seifert Van-Kampen $\pi_1(\mathcal{A} \cup \mathcal{B}) = \pi_1(\mathbb{R}^3 \setminus K)$ tiene la presentación deseada. ■

Notemos que la prueba anterior también vale para enlaces, eso da lugar al siguiente ejemplo.

3.6 Ejemplo. Si consideramos el *enlace de Hopf*, este tiene dos cruces y al orientarlo para obtener su presentación de Wirtinger como en la Figura 3.6, obtendremos que en ambos de los cruces se obtiene lo mismo y por lo tanto:

$$\pi_1(\mathbb{R}^3 \setminus \bigcirc \bigcirc) \cong \langle a_1, a_2 \mid a_2 a_1 a_2^{-1} a_1^{-1} \rangle,$$

sin embargo, esta es una presentación de $\mathbb{Z} \times \mathbb{Z}$, por lo tanto $\pi_1(\mathbb{R}^3 \setminus \bigcirc \bigcirc) \cong \mathbb{Z} \times \mathbb{Z}$. Directamente, podemos usar el teorema de Seifert-Van Kampen para mostrar que la unión disjunta de dos círculos sin ningún cruce tiene grupo fundamental $\langle a_1, a_2 \mid \emptyset \rangle$

el cual es isomorfo a $\mathbb{Z} * \mathbb{Z}$ por tanto estos dos enlaces no pueden ser equivalentes pues $\pi_1(\mathbb{R}^3 \setminus \bigcirc \bigcirc)$ es abeliano y $\pi_1(\mathbb{R}^3 \setminus \bigcirc \bigcirc)$ no. Esto muestra que, aunque ambos pueden verse como la unión disjunta de dos círculos, estos no son equivalentes.

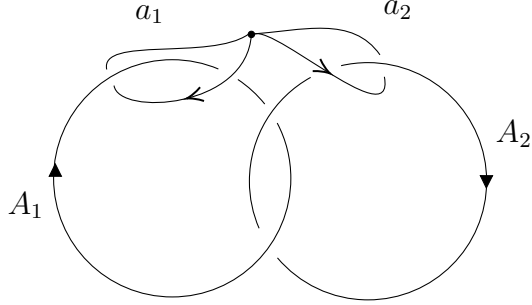


Figura 3.6. Diagrama del enlace de Hopf orientado.

3.7 Ejemplo. El nudo trébol es un nudo con tres cruces y por tanto su diagrama tendrá 3 componentes conexas, de esto tendremos 3 generadores y al bosquejarlo obtendremos que las relaciones vienen dadas por:

$$a_2 a_1^{-1} a_3^{-1} a_1 = 1, \quad (3.1)$$

$$a_3 a_2^{-1} a_1^{-1} a_2 = 1, \quad (3.2)$$

$$a_1 a_3^{-1} a_2^{-1} a_3 = 1. \quad (3.3)$$

Resolviendo para a_1 en la Ecuación 3.3 nos daremos que uno de los generadores es redundante y sustituyendo en las Ecuaciones 3.1 y 3.2 se obtiene en ambos casos:

$$a_3 a_2 a_3 = a_2 a_3 a_2. \quad (3.4)$$

De esto podremos obtener $(a_2 a_3 a_2)^2 = (a_3 a_2 a_3)(a_2 a_3 a_2) = (a_3 a_2)^3$ por lo que definiendo $a = a_2 a_3 a_2$, $b = a_3 a_2$, podremos notar que $a_2 = a b^2$ y $a_3 = b^2 a^{-1}$, eso nos da dos presentaciones equivalentes para el grupo de 3_1 , es decir:

$$\pi_1(\mathbb{R}^3 \setminus \bigcirc \bigcirc) \cong \langle a_1, a_2 \mid a_2 a_3 a_2 = a_3 a_2 a_3 \rangle \cong \langle a, b \mid a^2 = b^3 \rangle, \quad (3.5)$$

y con esta última equivalencia, podemos construir un homomorfismo sobreyectivo de $\pi_1(\mathbb{R}^3 \setminus \bigcirc \bigcirc)$ al grupo diédrico D_6 , por lo que el grupo del trébol no puede ser cíclico y por ende $\pi_1(\mathbb{R}^3 \setminus \bigcirc \bigcirc) \not\cong \pi_1(\mathbb{R}^3 \setminus \bigcirc)$. Esto muestra que el nudo trébol no puede ser equivalente al trivial y prueba la existencia de nudos no triviales.

3.2. Estados y brackets

Veremos también que al diagrama de un enlace le podemos asignar polinomios en función de los cruces de su diagrama. Para ello consideremos lo siguiente:

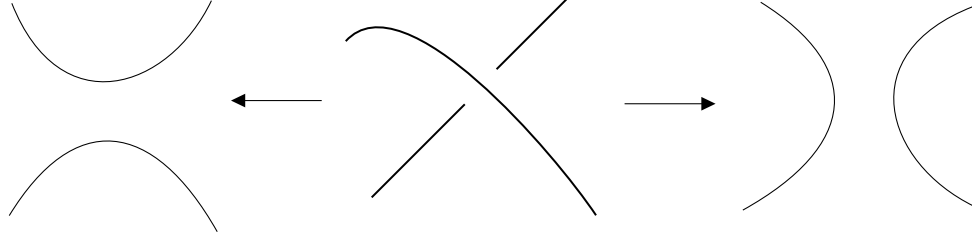


Figura 3.7. Descomposición de un cruce

Cada cruce $\times$ de un diagrama lo podemos *descomponer* de dos formas distintas, como en la Figura 3.7, y así mismo podemos fijarnos en que cada cruce separa localmente en cuatro regiones el ambiente que contiene el diagrama del enlace. De esta manera, identificando con las letras A y B dichas regiones en una vecindad del cruce y, descomponiendo todos los cruces obtendremos el *árbol de resoluciones para el diagrama de un enlace*. Con dicha identificación de las regiones, a las cuales llamaremos *etiquetas del estado*, podremos reconstruir el diagrama original del enlace como en la siguiente figura.

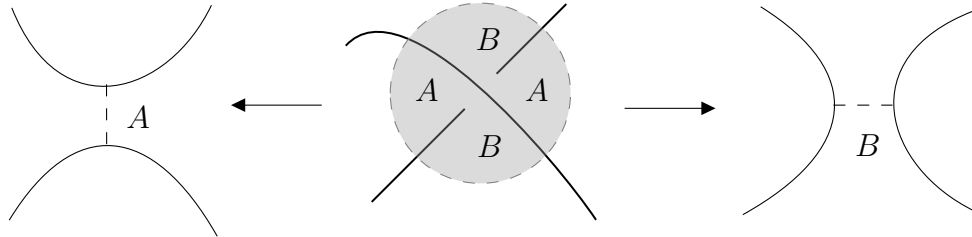


Figura 3.8. Descomposición e identificación de un cruce.

Cuando se dé una descomposición como la de la izquierda en la Figura 3.8, tendremos una *Etiqueta de tipo A* y cuando se dé una descomposición como la de la derecha tendremos una *Etiqueta de tipo B*. De esta manera obtendremos una familia de *descendientes del diagrama de un enlace*, entendiendo por descendientes del diagrama de un enlace a los diagramas obtenidos luego de deshacer uno de sus cruces de alguna de las dos maneras. Notemos que si iteramos este procedimiento con n cruces llegaremos a 2^n diagramas, cada uno compuesto por curvas de Jordan disjuntas.

3.8 Ejemplo. Si consideramos el nudo trébol 3_1 tendremos la siguiente descomposición.

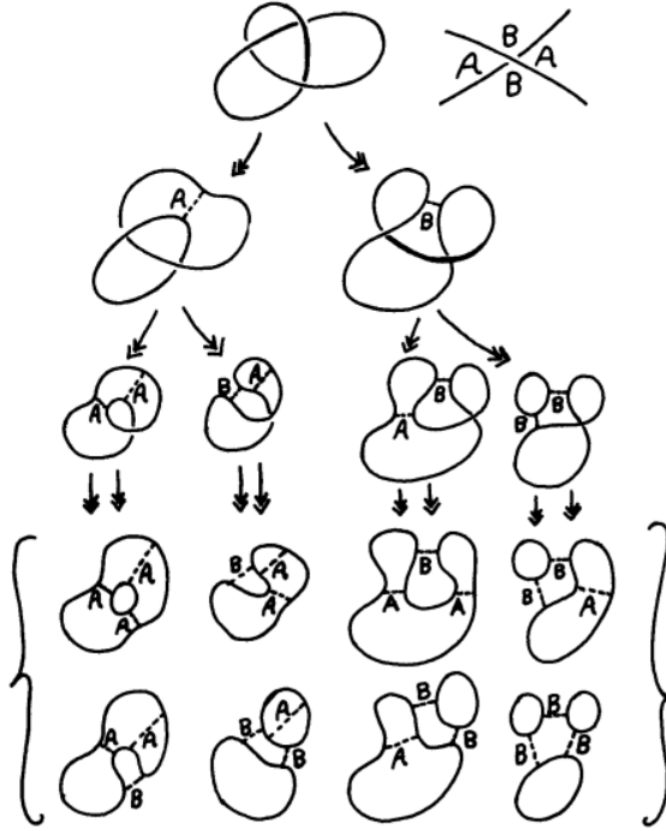


Figura 3.9. Descomposición del nudo trébol por Kauffman. Fuente: tomada de [18, p. 26].

A los últimos descendientes del diagrama L de un enlace los llamaremos *estados*, con los cuales definiremos las siguientes funciones:

3.9 Definición. Si L es diagrama de un enlace, $\mathcal{E}(L)$ es el conjunto de todos los estados de L y σ uno de sus estados, con σ_A el conjunto de las etiquetas de tipo A en el estado σ y σ_B el conjunto de las etiquetas de tipo B en el estado σ . El producto de etiquetas $\langle L|\cdot \rangle$ es tal que

$$\begin{aligned} \langle L|\cdot \rangle : \quad \mathcal{E}(L) &\rightarrow \mathbb{Z}[A, B], \\ \sigma &\mapsto \langle L|\sigma \rangle = A^{|\sigma_A|} B^{|\sigma_B|}, \end{aligned}$$

donde por $\mathbb{Z}[A, B]$ nos referimos al anillo de polinomios en las variables A y B con coeficientes en $\mathbb{Z}$.

3.10 Definición. Con la misma notación, si ℓ_σ es la cantidad de loops o curvas de Jordan en el estado σ entonces $\|\sigma\| := \ell_\sigma - 1$.

3.11 Ejemplo. Para uno de los estados del nudo trébol se tiene

$$\left\langle \text{trébol} \mid B \begin{array}{c} \circ \\ \circ \end{array} B \right\rangle = B^3, \quad \left\| \begin{array}{c} B \\ \circ \\ B \end{array} \right\| = 3 - 1 = 2$$

Con esto definimos finalmente el *polinomio bracket* o también conocido como el *bracket de Kauffman*:

3.12 Definición. Sea L el diagrama de un enlace con n cruces, entonces el *bracket de Kauffman* se define como el polinomio en tres variables $\langle L \rangle \in \mathbb{Z}[A, B, d]$ dado por:

$$\langle L \rangle = \langle L \rangle(A, B, d) = \sum_{i=1}^{2^n} \langle L | \sigma_i \rangle d^{\|\sigma_i\|}, \quad (3.6)$$

donde σ_i son los estados de L .

3.13 Ejemplo. Del árbol de descomposición del nudo trébol en la Figura 3.9, podremos obtener

$$\begin{aligned} \langle \text{trébol} \rangle &= A^3 d^{2-1} + A^2 B^1 d^{1-1} + A^2 B^1 d^{1-1} + A^1 B^2 d^{2-1} \\ &\quad + A^2 B^1 d^{1-1} + A^1 B^2 d^{2-1} + A^1 B^2 d^{2-1} + B^3 d^{3-1} \\ &= A^3 d + 3A^2 B + 3AB^2 d + B^3 d^2. \end{aligned}$$

Así, de la Definición 3.12 podemos deducir la siguientes propiedades:

3.14 Proposición. Si L es el diagrama de un enlace y $L \sqcup \bigcirc$ denota la unión disjunta de L (sin ningún cruce) con una copia del nudo trivial,

- a) $\langle \times \rangle = A \langle \nearrow \rangle + B \langle \searrow \rangle$
- b) $\langle L \sqcup \bigcirc \rangle = d \langle L \rangle$
- c) $\langle \bowtie \rangle = AB \langle \rangle + (A^2 + B^2 + ABd) \langle \nearrow \rangle$
- d) $\langle \text{trébol} \rangle = (Ad + B) \langle \text{trébol} \rangle$
- e) $\langle \text{trébol} \rangle = (A + Bd) \langle \text{trébol} \rangle$

Así como lo definimos, el bracket de Kauffman no es un invariante en nuestro contexto, pues las propiedades c) y d) exhiben que en general no tendremos el mismo bracket Kauffman si aplicamos el primer movimiento de Reidemeister a nuestro diagrama de enlace, es decir, el bracket de Kauffman no es invariante bajo isotopías de ambiente. Sin embargo, podemos buscar condiciones para que el bracket sea invariante en nuestros términos, es decir, que dos enlaces isotópicos tengan el mismo polinomio, esto imponiendo condiciones sobre las variables A, B y d .

Ahora, para que el bracket de Kauffman sea invariante bajo isotopías regulares, únicamente debemos probar que es invariante bajo los movimientos de Reidemeister 2 y 3. En otras palabras, lo que estamos pidiendo es que

$$\langle \rangle = \langle \bowtie \rangle = AB \langle \rangle + (A^2 + B^2 + ABd) \langle \smile \rangle.$$

Una forma de lograr esta igualdad sería la siguiente

$$\begin{aligned} AB &= 1 \\ A^2 + B^2 + ABd &= 0. \end{aligned}$$

Con esto, obtenemos que una condición suficiente para tener un invariante bajo $\mathcal{R}_2$ es $B = A^{-1}$ y $d = -A^2 - A^{-2}$ y podemos probar lo siguiente:

3.15 Proposición. *El bracket de Kauffman con la sustitución $B = A^{-1}$ y $d = -A^2 - A^{-2}$ del diagrama de un enlace L es invariante bajo el movimiento de Reidemeister 3. Esto significa que si L' es obtenido del diagrama de un enlace L por una iteración de movimientos de Reidemeister 3 entonces*

$$\langle L \rangle(A, A^{-1}, -A^2 - A^{-2}) = \langle L' \rangle(A, A^{-1}, -A^2 - A^{-2})$$

Demostración. Dado que aplicar un movimiento de Reidemeister es un proceso local, para la prueba basta con fijarnos en la parte en la que se realiza $\mathcal{R}_3$. Luego usando la invarianza bajo $\mathcal{R}_2$ se tiene:

$$\begin{aligned} \langle \text{diagrama 3.1} \rangle &= A \langle \text{diagrama 3.2} \rangle + B \langle \text{diagrama 3.3} \rangle \\ &= A \langle \text{diagrama 3.4} \rangle + B \langle \text{diagrama 3.5} \rangle \\ &= A \langle \text{diagrama 3.6} \rangle + B \langle \text{diagrama 3.7} \rangle \\ &= \langle \text{diagrama 3.8} \rangle \end{aligned}$$

■

Por ello, desde este momento nos referiremos a $\langle L \rangle(A, A^{-1}, -A^2 - A^{-2})$, simplemente, como $\langle L \rangle$.

Nuevamente, podremos darnos cuenta que, por las propiedades $d)$ y $e)$ en la Proposición 3.14, este no es un invariante bajo $\mathcal{R}_1$ pues:

$$\langle \text{cruce} \rangle = (A(-A^2 - A^{-2}) + A^{-1})\langle \text{curva} \rangle = (-A^3)\langle \text{curva} \rangle, \quad (3.7)$$

$$\langle \text{cruce} \rangle = (A + A^{-1}(-A^2 - A^{-2}))\langle \text{curva} \rangle = (-A^{-3})\langle \text{curva} \rangle. \quad (3.8)$$

Ahora, para obtener por completo un invariante bajo isotopías de ambiente, modificaremos un poco el bracket de Kauffman.

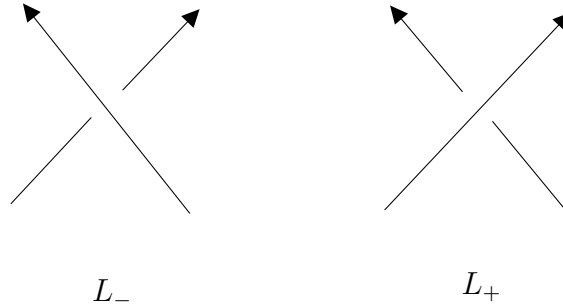


Figura 3.10. Cruces con orientación.

Empezaremos dotando de una orientación a un nudo y luego asignándole un signo a cada cruce, como en la Figura 3.10. Donde si $\epsilon \in \{1, -1\}$ tendremos que $\epsilon(L_-) = -1$ y $\epsilon(L_+) = +1$. Con esto definimos la *torcedura*, también conocida como *writhe* en inglés.

3.16 Definición. Sea L un diagrama de un enlace con una orientación. Definimos la *torcedura* de L , como

$$w(L) := \sum_p \epsilon(p), \quad (3.9)$$

donde p recorre todos los cruces del diagrama L .

Con esto obtendremos nuestro candidato a invariante bajo los tres movimientos de Reidemeister.

3.17 Definición. Definimos el *bracket normalizado* $\mathcal{L}_L$ para L , el diagrama de un

enlace como

$$\mathcal{L}_L = (-A^3)^{-w(L)} \langle L \rangle. \quad (3.10)$$

Luego, de las ecuaciones 3.7 y 3.8 se sigue este si es invariante bajo $\mathcal{R}_1$. Como $w(L)$ es invariante bajo isotopías regulares obtendremos que $\mathcal{L}_L$ es un invariante bajo $\mathcal{R}_1$, $\mathcal{R}_2$ y $\mathcal{R}_3$, es decir, invariante bajo isotopías de ambiente.

3.3. Polinomios de Jones y de Alexander-Conway

Antes que nada, notemos que una expresión de la forma $t + t^{-1}$ realmente no es un polinomio en el sentido estricto de la definición. Sin embargo, existe una noción más general de polinomio la cual será la que manejaremos de ahora en adelante.

3.18 Definición. Dados $m, n \in \mathbb{Z}$ con $n \leq m$, a la siguiente expresión

$$p(x) = \sum_{n \leq k \leq m} a_k x^k \quad (3.11)$$

la llamaremos *polinomio de Laurent en x* . Nos referiremos al conjunto de todos los polinomios de Laurent donde $a_k \in \mathbb{Z}$ como $\mathbb{Z}[x, x^{-1}]$.

Notemos que con esta definición, el bracket de Kauffman y el bracket normalizado son polinomios de Laurent. Ahora con esto definiremos uno de los invariantes más famosos de la teoría de nudos, el polinomio de Jones.

3.19 Definición. El *polinomio de Jones* $V_L(t)$ de un enlace orientado L , es un polinomio de Laurent en la variable $t^{1/2}$ que se define de la siguiente forma

1. Si existe una isotopía de ambiente de L a L' , entonces $V_L(t) = V_{L'}(t)$.
2. $V_{\bigcirc}(t) = 1$
3. $t^{-1}V_{\nearrow\searrow}(t) - tV_{\searrow\swarrow}(t) = \left(\sqrt{t} - \frac{1}{\sqrt{t}}\right) V_{\swarrow\nwarrow}$

Con esta definición se puede calcular el polinomio de Jones para cualquier nudo. Sin embargo, podemos calcularlo a través del bracket normalizado teniendo en cuenta el siguiente resultado:

3.20 Teorema. Sea $\mathcal{L}_L(A) = (-A^3)^{-w(L)} \langle L \rangle$, entonces

$$\mathcal{L}_L(t^{-1/4}) = V_L(t). \quad (3.12)$$

La prueba de esto puede ser encontrada en [18, p. 50]. De manera similar es posible definir el polinomio de Conway-Alexander.

3.21 Definición. Si es z una variable, el polinomio de Alexander $\nabla(z)$ (O Conway-Alexander) de un enlace orientado L , es un polinomio en $\mathbb{Z}[z]$ que se define de la siguiente forma

1. Si existe una isotopía de ambiente de L a L' , entonces $\nabla_L(z) = \nabla_{L'}(z)$.
2. $\nabla_{\bigcirc}(z) = 1$
3. $\nabla_{\times}(z) - \nabla_{\searrow}(z) = z\nabla_{\nearrow}(z)$

3.4. Homología de Khovanov

Para esta sección debemos volver a hablar un poco del corchete de Kauffman y enfocarnos en algunos de los detalles de esta construcción para estandarizar la notación y pasar del enfoque hecho por Kauffman en [18] a uno más moderno. En particular, parte del panorama actual de este trabajo se encuentra recopilado en [25]. Si tomamos la Definición 3.12 y haciendo la sustitución anteriormente discutida obtendremos la siguiente expresión para el bracket de Kauffman $\langle L \rangle$:

$$\langle L \rangle = \sum_{\sigma \in \mathcal{E}(L)} A^{|\sigma_A| - |\sigma_B|} (-A^2 - A^{-2})^{||\sigma||}. \quad (3.13)$$

A este polinomio de Laurent a veces se le llama *polinomio bracket reducido*. Con la ecuación 3.13 o directamente del Ejemplo 3.13, el polinomio bracket reducido del trébol se simplifica a $A^{-7} - A^5 - A^3$. Al conjunto de todos los loops o curvas de Jordan de un estado σ lo denotaremos por D_σ y por ende $|D_\sigma|$ es el número de curvas de Jordan en el estado σ . Con esta nueva notación podemos notar que $||\sigma|| = |D_\sigma| - 1$ y $\ell_\sigma = |D_\sigma|$. Ahora presentaremos otra versión del polinomio bracket.

3.22 Definición. El *polinomio bracket no reducido* se define para L el diagrama de un enlace, como $[\emptyset] := 1$, $[\bigcirc] := -A^2 - A^{-2}$ y en general:

$$[L] := (-A^2 - A^{-2})\langle L \rangle = \sum_{\sigma \in \mathcal{E}(L)} A^{|\sigma_A| - |\sigma_B|} (-A^2 - A^{-2})^{||\sigma||+1}. \quad (3.14)$$

Con esta definición tendremos:

$$[\text{trébol}] = -A^{-9} + A^7 + A^3 + A^{-1}. \quad (3.15)$$

3.23 Definición. Un *estado de Kauffman mejorado* $S(\sigma, \varepsilon)$ para el diagrama de un enlace L , es un estado σ junto a una función $\varepsilon : D_\sigma \rightarrow \{+, -\}$, la cual asigna a cada curva de Jordan en D_σ un signo.

Al conjunto de los estados mejorados del diagrama de un enlace L lo llamaremos $\mathcal{EKS}(L)$ por ser las siglas de su nombre en inglés, Enhanced Kauffman States. Dado que cada curva de Jordan en D_σ tiene dos posibilidades (es decir dos posibles imágenes bajo ε), para cada estado σ se tienen $2^{||\sigma||+1} = 2^{|D_\sigma|}$ estados mejorados. De esta definición, podemos expresar el polinomio bracket no reducido en función de monomios.

3.24 Proposición. Para L el diagrama de un enlace, su polinomio bracket no reducido $[L]$ viene dado por:

$$[L] = \sum_{S(\sigma, \varepsilon) \in \mathcal{EKS}(L)} (-1)^{|D_\sigma|} A^{|\sigma_A| - |\sigma_B|} (A^2)^{|\varepsilon^{-1}(+)| - |\varepsilon^{-1}(-)|}. \quad (3.16)$$

Demostración. Lo primero que podemos notar es que para todo $\sigma \in \mathcal{E}(L)$, si $E(\sigma)$ el conjunto de todos los posibles estados mejorados para σ , entonces:

$$(A^2 + A^{-2})^{|D_\sigma|} = \sum_{S(\sigma, \varepsilon) \in E(\sigma)} (A^2)^{|\varepsilon^{-1}(+)| - |\varepsilon^{-1}(-)|}. \quad (3.17)$$

Esto se debe a que al expandir la expresión

$$\underbrace{(A^2 + (A^{-1})^2) \cdots (A^2 + (A^{-1})^2)}_{|D_\sigma| \text{ veces}},$$

sin agrupar, obtendremos $2^{|D_\sigma|}$ términos que corresponden a las posibles tuplas de la forma $(A^{\pm 2}, A^{\pm 2}, \dots, A^{\pm 2})$. Por otro lado cada tupla corresponde a un elemento de $E(\sigma)$, pues cada posición en la tupla se puede identificar con una curva de Jordan. Debido a esta biyección:

$$\begin{aligned} (A^2 + A^{-2})^{|D_\sigma|} &= \sum_{k=1}^{2^{|D_\sigma|}} (A^2)^{|\varepsilon_k^{-1}(+)| - |\varepsilon_k^{-1}(-)|} \\ &= \sum_{S(\sigma, \varepsilon) \in E(\sigma)} (A^2)^{|\varepsilon^{-1}(+)| - |\varepsilon^{-1}(-)|}. \end{aligned}$$

Por lo que usando la ecuación 3.17 tendremos

$$\begin{aligned}
[L] &= \sum_{\sigma \in \mathcal{E}(L)} A^{|\sigma_A| - |\sigma_B|} (-A^2 - A^{-2})^{|D_\sigma|}, \\
&= \sum_{\sigma \in \mathcal{E}(L)} (-1)^{|D_\sigma|} A^{|\sigma_A| - |\sigma_B|} (A^2 + A^{-2})^{|D_\sigma|}, \\
&= \sum_{\sigma \in \mathcal{E}(L)} (-1)^{|D_\sigma|} A^{|\sigma_A| - |\sigma_B|} \sum_{S(\sigma, \varepsilon) \in E(\sigma)} (A^2)^{|\varepsilon^{-1}(+)| - |\varepsilon^{-1}(-)|}, \\
&= \sum_{\sigma \in \mathcal{E}(L)} \sum_{S(\sigma, \varepsilon) \in E(\sigma)} (-1)^{|D_\sigma|} A^{|\sigma_A| - |\sigma_B|} (A^2)^{|\varepsilon^{-1}(+)| - |\varepsilon^{-1}(-)|}, \\
&= \sum_{S(\sigma, \varepsilon) \in \mathcal{EK}\mathcal{S}(L)} (-1)^{|D_\sigma|} A^{|\sigma_A| - |\sigma_B|} (A^2)^{|\varepsilon^{-1}(+)| - |\varepsilon^{-1}(-)|}.
\end{aligned}$$

■

De esta manera conviene fijarnos en los términos $|\sigma_A| - |\sigma_B|$ y $|\varepsilon^{-1}(+)| - |\varepsilon^{-1}(-)|$. Por ende denotamos por $\varsigma(\sigma)$ a la diferencia entre el numero de etiquetas de tipo A (también llamados marcadores positivos) y etiquetas de tipo B (marcadores negativos) del estado σ . Así mismo se denota por $\tau(S(\sigma, \varepsilon))$ a la diferencia entre el numero de signos positivos y negativos en el estado mejorado, esto es:

$$\varsigma(\sigma) = |\sigma_A| - |\sigma_B|, \quad \tau(S(\sigma, \varepsilon)) = |\varepsilon^{-1}(+)| - |\varepsilon^{-1}(-)|.$$

Luego, podremos modificar la Ecuación 3.16 y obtener:

$$[L] = \sum_{S(\sigma, \varepsilon) \in \mathcal{EK}\mathcal{S}(L)} (-1)^{|D_\sigma|} A^{\varsigma(\sigma) + 2\tau(S(\sigma, \varepsilon))}, \quad (3.18)$$

la cual se conoce como la formula para el polinomio bracket no reducido en términos de los estados mejorados. Estos estados mejorados forman una base para los grupos de cadena y el complejo de cadena de Khovanov.

3.25 Definición. Por *bigrado* en los estados mejorados nos referiremos al conjunto

$$\mathcal{S}_{a,b}(L) = \mathcal{S}_{a,b} = \{S(\sigma, \varepsilon) \in \mathcal{EK}\mathcal{S}(L) : a = \varsigma(\sigma), b = \varsigma(\sigma) + 2\tau(S(\sigma, \varepsilon))\}.$$

Los *grupos de cadena* $\mathcal{C}_{a,b}(L) = \mathcal{C}_{a,b}$ se definen como los grupos abelianos libres con base $\mathcal{S}_{a,b}(L) = \mathcal{S}_{a,b}$, esto es $\mathcal{C}_{a,b} = \mathbb{Z}\langle \mathcal{S}_{a,b} \rangle$. Con esto

$$\mathcal{C}(L) = \bigoplus_{a,b \in \mathbb{Z}} \mathcal{C}_{a,b},$$

es el *grupo abeliano libre bigradado*.

3.26 Definición. Para L el diagrama de un enlace, definimos *el complejo de cadena* $\mathcal{C}(L) = \{(\mathcal{C}_{a,b}, \partial_{a,b})\}$, donde $\partial_{a,b} : \mathcal{C}_{a,b} \rightarrow \mathcal{C}_{a-2,b}$ es el *operador frontera* definido por:

$$\partial_{a,b}(S) = \sum_{S' \in \mathcal{S}_{a-2,b}} (-1)^{t(S,S')} (S, S') S'.$$

Donde (S, S') es lo que se conoce como el *número incidental* entre S y S' . Tenemos que $(S, S') \in \{0, 1\}$ y $(S, S') = 1$ si y solamente si:

1. S y S' son idénticos, excepto en un cruce, denotado por v , además la etiqueta de v en el estado σ de S es de tipo A y la etiqueta de v en el estado σ' de S' es de tipo B.
2. $\tau(S') = \tau(S) + 1$ y cada componente de D_σ que no interactuá con el cruce v mantiene su signo por $D_{\sigma'}$.

Además, dado un orden en los cruces en el diagrama L , es número $t(S, S')$ es definido como el número de cruces en S con etiqueta de tipo B después del cruce v en el orden dado.

El orden del que se habla en la última parte de esta definición induce la relación $\partial_{a-2,b} \circ \partial_{a,b} = 0$ y $\text{im}(\partial_{a+2,b}) \subseteq \text{Ker}(\partial_{a,b})$ lo que permite definir los siguientes grupos.

3.27 Definición. La *homología de Khovanov* para L el diagrama de un enlace, se define como la homología del complejo de cadena $\mathcal{C}(L)$:

$$H_{a,b}(L) = \frac{\text{Ker}(\partial_{a,b})}{\text{im}(\partial_{a+2,b})}.$$

Siendo, como podría esperar el lector, la propiedad más importante la invarianza de $H_{a,b}$ bajo isotopías regulares. Para los detalles sobre la demostración de este resultado y la relación entre $H_{a,b}(L_1)$ y $H_{c,d}(L_2)$ cuando L_1 y L_2 están relacionados por un movimiento de Reidemeister del primer tipo $\mathcal{R}_1$, puede consultarse [25],[39] y [38]. Para una introducción al concepto general de homología y algunas definiciones auxiliares para las construcciones de este capítulo, se recomienda consultar el Apéndice B.

4. CONEXIONES CON FÍSICA

Muchas veces ocurre que el trabajo interdisciplinario da como resultado el enriquecimiento de las áreas que se ven involucradas. Teniendo esto en mente y cambiando un poco el tono abstracto de los anteriores capítulos, discutiremos, de manera un poco informal, algunos modelos físicos que están profundamente relacionados con la teoría de nudos, esto basándonos en el trabajo de Kauffman y algunas de sus interpretaciones en [18], pues es conocido por sus numerosos aportes a la teoría de nudos. Muchos de estos resultados vendrán de distintas áreas de la física como la física de estado sólido, o la física de partículas.

4.1. Funciones de Partición

Para iniciar, nos centraremos un poco en mecánica estadística, y veremos que la construcción del bracket de Kauffman tiene una interpretación curiosa, la cual el mismo Kauffman menciona en [18].

Un escenario muy común es tener un sistema A que consiste en un número fijo N de partículas en un volumen V , pero la única información disponible de la energía del sistema es un promedio de la energía $\overline{E}$, donde para encontrar esta media debemos calcular la energía de todos los posibles estados de este sistema y promediarlos. Dicho de otra manera, si el sistema A tiene a estados y si la energía de cada estado es $E(r)$, donde cada estado se repite a_r veces entonces:

$$\frac{1}{a} \sum_r a_r E(r) = \overline{E}. \quad (4.1)$$

A cada uno de los estados con energía $E(r)$ se les conoce como microestados. Se denomina ensamble canónico, al conjunto de los posibles estados de un sistema (conjunto de partículas) que intercambia energía térmica con los alrededores, pero no materia.

Los sistemas en su ensamble estadístico representativo se distribuyen sobre todos los estados accesibles de acuerdo a la distribución canónica. Para estas situaciones obtendremos que el promedio de la energía es

$$\overline{E} = \frac{\sum_r e^{-\beta E(r)} E(r)}{\sum_r e^{-\beta E(r)}}, \quad (4.2)$$

donde el parámetro β se escoge precisamente para tener esta igualdad. Notemos ahora que la expresión del numerador puede escribirse como:

$$\sum_r e^{-\beta E(r)} E(r) = - \sum_r \frac{\partial}{\partial \beta} (e^{-\beta E(r)}) = - \frac{\partial}{\partial \beta} \left(\sum_r e^{-\beta E(r)} \right), \quad (4.3)$$

de esta manera podemos reescribir 4.2 como:

$$\overline{E} = - \frac{\partial}{\partial \beta} \left(\ln \left[\sum_r e^{-\beta E(r)} \right] \right), \quad (4.4)$$

por lo que si

$$\mathcal{Z}(\beta) := \sum_r e^{-\beta E(r)}, \quad (4.5)$$

vale la pena fijarnos en la función $\mathcal{Z}(\beta)$, la cual recibe un nombre, se le conoce como *función de partición*. Más específicamente, el parámetro β se toma como $\beta = \frac{1}{k_B T}$ donde k_B es la constante de Boltzmann y T es temperatura del sistema. En general, una función de partición puede definirse de la siguiente forma:

4.1 Definición. Dado un conjunto de variables aleatorias X_i tomando valores x_i y algún tipo de función potencial o Hamiltoniano $H(x_1, x_2, \dots)$ la función de partición $\mathcal{Z}(\beta)$ se define como

$$\mathcal{Z}(\beta) = \sum_{x_i} e^{-\beta H(x_1, x_2, \dots)}, \quad (4.6)$$

donde β es un parámetro real.

De entrada podemos ver que estas funciones guardan muchas similitudes con la Definición 3.12, más concretamente si reescribimos la Ecuación 3.6 de la forma:

$$\langle L \rangle = \sum_{\sigma \in \mathcal{E}(L)} \langle L | \sigma \rangle d^{||\sigma||}, \quad (4.7)$$

esta se asemeja bastante a la ecuación expuesta en [3, p. 9] para $\langle X \rangle$, el valor termodinámico observado de una propiedad X con valor $X(\sigma)$ para el estado σ de un sistema:

$$\langle X \rangle = \sum_{\sigma} \frac{X(\sigma)}{\mathcal{Z}(1/k_B T)} (e^{-1/k_B T})^{E(\sigma)}. \quad (4.8)$$

Por ello, si vemos al diagrama de un nudo como un sistema y a cada uno de los últimos descendientes como un posible estado, el bracket de Kauffman puede interpretarse como una función de partición pues está sumando sobre todos estados y es invariante bajo ciertas transformaciones del sistema. Siendo esta una interpretación que Kauffman usa directamente en [18] para enfatizar la importancia de la Ecuación 4.7. Sin embargo, hay una conexión aún más directa entre las funciones de partición y los nudos, la cual parte de la siguiente definición:

4.2 Definición. Del polinomio en tres variables de la Ecuación 4.7, al hacer la sustitución $A = q^{-\frac{1}{2}}v$, $B = 1$ y $d = q^{\frac{1}{2}}$, se obtiene el *bracket especial*:

$$\{L\} = \{L\}(q, v) = \langle L \rangle \left(q^{-\frac{1}{2}}v, 1, q^{\frac{1}{2}} \right). \quad (4.9)$$

Por otro lado, a cada grafo planar G se le puede asignar, el diagrama de un enlace $K(G)$, mediante un proceso brevemente descrito en [17, pp. 25-27]. Más aún, es posible probar una correspondencia uno a uno entre diagramas de enlaces y cierto tipo de grafos planares, los detalles de dicha correspondencia y una demostración de esto pueden ser encontradas en [16]. Así, otro objeto sumamente relacionado con los grafos planares es el siguiente:

4.3 Definición. El polinomio dicromático $Z_G = Z_G(q, v)$ de un grafo plano G , definido por las siguientes fórmulas:

$$\begin{aligned} Z_G &= Z_{G'} + vZ_{G''}, \\ Z_{G \sqcup \bullet} &= qZ_G, \\ Z_{\bullet} &= q. \end{aligned}$$

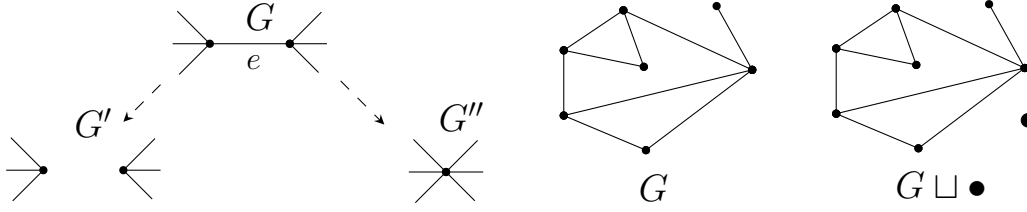


Figura 4.1. Procedimientos para el cálculo de un polinomio dicromático.

Donde G' es el grafo obtenido borrando una arista e del grafo G , G'' es obtenido mediante la contracción de los dos vértices que une la arista e en un solo vértice, $G \sqcup \bullet$ es la unión disjunta de G con un vértice (es decir, el grafo agregando un vértice sin ninguna arista que lo una al grafo original) y $Z_\bullet$ es el caso para un solo vértice aislado, todo esto descrito en la Figura 4.1.

Con esto, un resultado de Kauffman mostrado en [15] es el siguiente:

4.4 Proposición. *El polinomio dicromático Z_G de un grafo planar G está relacionado con el bracket especial mediante la ecuación:*

$$Z_G = q^{\frac{N}{2}} \{K(G)\}, \quad (4.10)$$

donde N es el número de vértices de G .

Ya con estas definiciones aclaradas y la proposición 4.4 expuesta, la conexión que nos interesa esta vez está dada por el polinomio dicromático, pues al hacer la sustitución $v = e^{-1/k_B T} - 1$, Temperley y Lieb probaron en [37] que $Z_G(q, v)$ es la función de partición del *Modelo Potts*, muy importante en física de estado sólido. De esta manera, por la ecuación 4.10 tenemos que un caso particular del bracket de Kauffman, no solo puede ser interpretado como una función de partición, es directamente, una función de partición.

4.2. Tensores abstractos y la ecuación de Yang-Baxter

Existe un formalismo en el que un tensor puede ser representado gráficamente para, posteriormente, realizar operaciones entre estos diagramas, para dicha representación, consideremos un gráfico como el de la Figura 4.2, el cual será el que

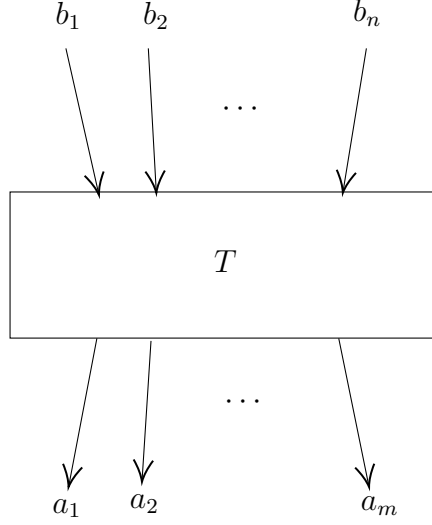


Figura 4.2. Tensor abstracto $T_{a_1 \dots a_m}^{b_1 \dots b_n}$.

representa un *tensor abstracto* con n índices superiores y m índices inferiores, escribiendo en este caso: $T_{a_1 \dots a_m}^{b_1 \dots b_n}$.

Al estar familiarizados con las convenciones para realizar operar tensores en términos de símbolos, esta notación tendrá ciertas ventajas que discutiremos más adelante. Debemos resaltar que Kauffman en [18, p. 104] decide nombrar a estos objetos como *tensores abstractos diagramáticos* pues, realmente, lo que tenemos es una versión visual para el álgebra matricial, cuando una matriz tiene una mayor cantidad de índices. Por esto y porque estos diagramas pueden ser transformados mediante ciertas propiedades es que también podemos llamarlos tensores abstractos. Para representar una matriz $M = (M_j^i)$ con esta notación, podemos dibujar un

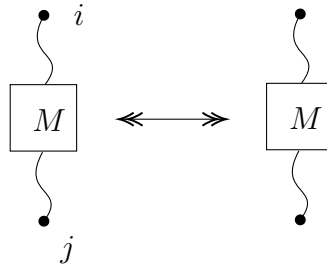


Figura 4.3. Diagrama como una matriz.

cuadrado con dos segmentos unidos, como en la Figura 4.3, donde en este caso los índices superiores e inferiores de M_j^i estarán en el segmento superior y los índices inferiores estarán en el segmento inferior. De esta manera, al usar la notación de Einstein para la traza de una matriz $M = (M_j^i)$ y el producto de dos matrices

$$A = (A_j^i), B = (B_j^i)$$

$$(AB)_j^i = \sum_k A_k^i B_j^k = A_k^i B_j^k,$$

$$\text{tr}(M) = \sum_i M_i^i = M_i^i.$$

Podremos expresar términos de diagramas estas operaciones, tal y como en la Figura 4.4. Lo cual es una ventaja pues nos muestra que los índices pueden ser omitidos pues la posición de los segmentos dentro del diagrama rescata esa información.

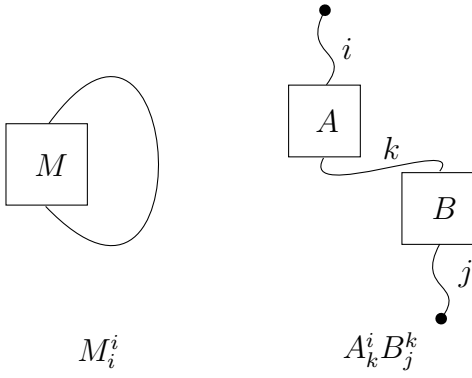


Figura 4.4. Operaciones de Matrices.

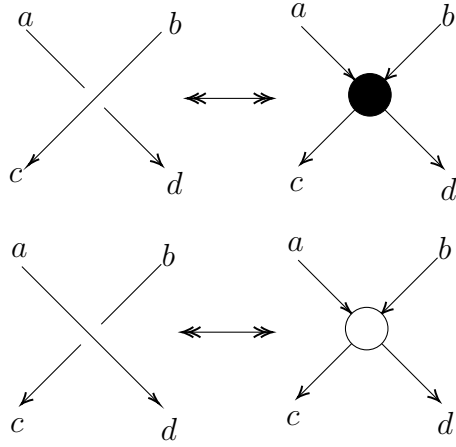


Figura 4.5. Cruces identificados.

Ahora, haremos un proceso parcialmente inverso, tomaremos el diagrama de un nudo o un enlace y lo representaremos como un tensor abstracto, usando ideas muy parecidas a las discutidas en esta sección. Escogiendo una orientación para el enlace, los dos tipos posibles de cruces dan a lugar a tensores abstractos, tal y como en la Figura 4.5, en concreto, se tiene $T(\times) = R_{cd}^{ab}$ y $T(\oslash) = \bar{R}_{cd}^{ab}$. Esto codifica la información de un enlace L en un tensor abstracto $T(L)$.

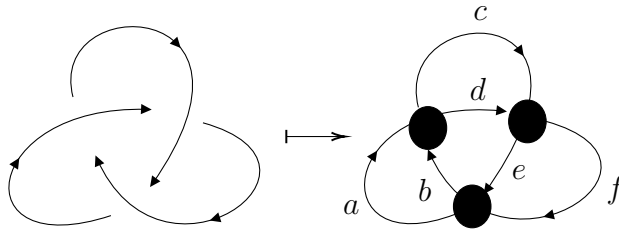


Figura 4.6. Trébol orientado con cruces identificados.

4.5 Ejemplo. El tensor asociado al nudo trébol 3_1 , orientado como en la Figura 4.6, es el siguiente:

$$T(\text{trébol } 3_1) = R_{dc}^{ba} R_{ef}^{dc} R_{ba}^{ef}. \quad (4.11)$$

Con esta notación, podemos expresar la condición para la cual el tensor $T(L)$ asociado al enlace L , es invariante bajo el tercer movimiento de Reidemeister, vistas en la Figura 4.7.

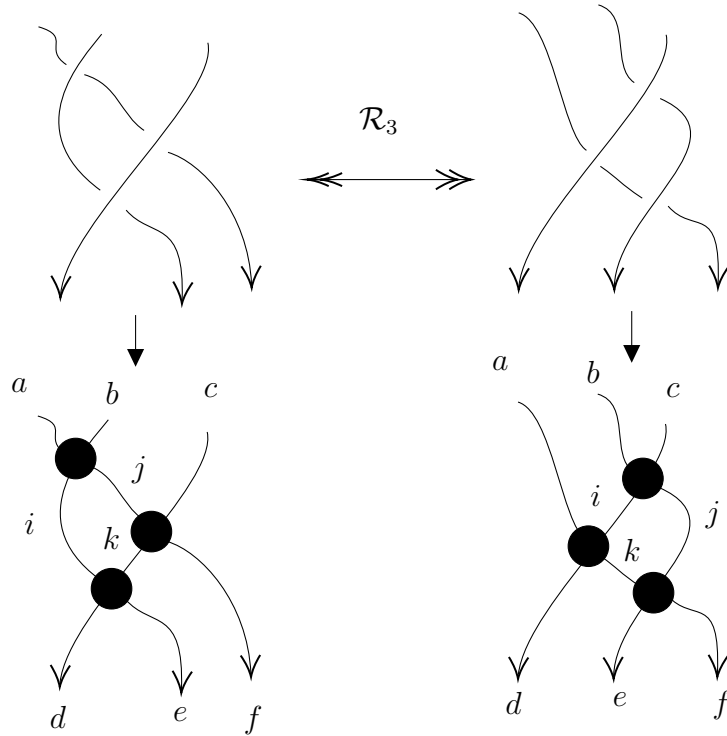


Figura 4.7. Movimiento de Reidemeister 3 para un enlace orientado en el que todos sus cruces son positivos y están identificados.

De esta manera, la condición de invarianza para el caso de tener solo cruces positivos, viene dada por la ecuación de *Yang Baxter*:

$$\sum_{i,j,k} R_{ij}^{ab} R_{kf}^{jc} R_{de}^{ik} = \sum_{i,j,k} R_{ij}^{bc} R_{dk}^{ai} R_{ef}^{kj}. \quad (4.12)$$

La cual al emplear la notación de Einstein es:

$$R_{ij}^{ab} R_{kf}^{jc} R_{de}^{ik} = R_{ij}^{bc} R_{dk}^{ai} R_{ef}^{kj}. \quad (4.13)$$

La ecuación cuando hay otros tipos de cruces se construye de manera análoga y

siempre guarda ciertas similitudes con la Ecuación 4.13. Esta formulación se debe Kauffman y es descrita con más detalle en [18, pp. 104-110]. Además, existe una interpretación para los diagramas descritos en la Figura 4.5, pues estos vértices pueden verse como interacciones entre partículas, donde los índices a, b, c, d son ciertas cantidades asociadas a las partículas. Esta interpretación, nos lleva a pensar en diagramas de Feynman.

Precisamente, es posible establecer una conexión entre teoría de nudos y la evaluación de diagramas de Feynman, sin embargo, para ello deben considerarse los siguientes diagramas y asociarlos a tensores: habrán 5 tensores en los que debemos centrarnos, a los cuales llamaremos *cuencas* M^{ab} , *cúspides* M_{ab} , *arcos* δ_a^b , *cruce uno* R_{cd}^{ab} y *cruce dos* $\bar{R}_{cd}^{ab}$, como en la Figura 4.8. En el caso de las cuencas pueden ser definidas, formalmente, como un mínimo local, las cúspides pueden ser definidas, más formalmente como máximos locales y los arcos, simplemente, como una curva en la que no se alcanza ni un máximo ni un mínimo.

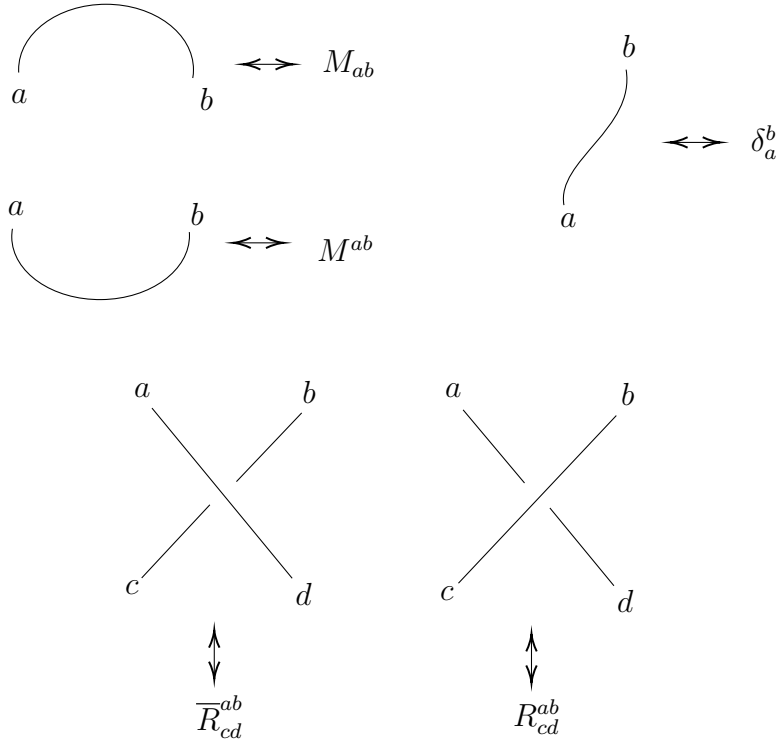


Figura 4.8. Tensores abstractos utilizados.

Los tensores M_{ab} , M^{ab} , δ_a^b , R_{cd}^{ab} y $\bar{R}_{cd}^{ab}$, módulo ciertas relaciones descritas en [23], ayudan a dar una interpretación expuesta en [18, pp. 117-147] como diagramas de Feynman.

4.3. Teorema de la estadística del espín

La última conexión surge de una dificultad con la que nos topamos con anterioridad. En las Ecuaciones 3.7 y 3.8, nos fijamos que el bracket de Kauffman no es invariante bajo movimientos $\mathcal{R}_1$, esto se debe a que en realidad esto no siempre es una “equivalencia topológica” pues si en lugar de considerar una cuerda unidimensional la consideramos como superficie, veremos que deshacer uno de estos cruces realmente nos lleva a tener giros en la misma superficie (claro si dejamos fijo el principio y el final de lo que antes era la cuerda).

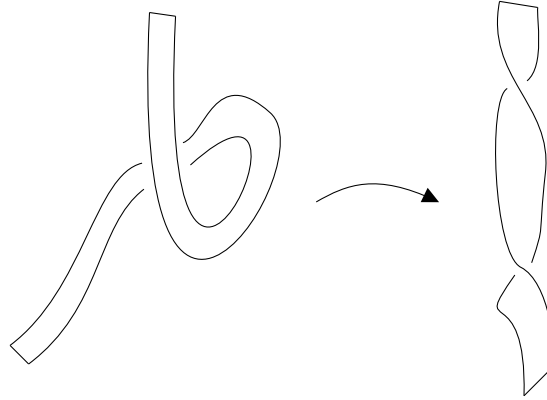


Figura 4.9. Movimiento $\mathcal{R}_1$ aplicado a un cruce de una superficie.

Esto recuerda a la alegoría de Dirac para el espín, la cual ilustra el modelo usando cuaterniones detrás del espín. Esta alegoría una de las muchas formas de explicar que, el espín de ciertas partículas vuelve a su estado original solo después de dos rotaciones completas, no después de una. También puede verse como un idea intuitiva del porque el grupo de matrices complejas,

$$SU(2) = \left\{ \begin{pmatrix} \alpha & -\bar{\beta} \\ \beta & \bar{\alpha} \end{pmatrix} : \alpha, \beta \in \mathbb{C}, \quad |\alpha|^2 + |\beta|^2 = 1 \right\}, \quad (4.14)$$

es simplemente conexo.

Aunque la naturaleza de este argumento es heurística e ilustrativa, es citada incluso por Kauffman en [18, pp. 403-426]. También se utiliza una de sus variantes para explicar el *Teorema de la Estadística del Espín*, el cual caracteriza partículas como bosones y fermiones mediante espines. En [33] se habla más a detalle sobre esto, pero la idea principal es que por la Ecuación 3.7, podemos escribir el deshacer un

cruce en una superficie parecida a un cinturón como en la Figura 4.10 e interpretar este $-A^3$ como un factor de rotación.

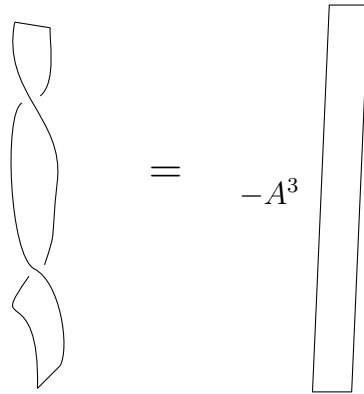


Figura 4.10. Expresando cómo deshacer un cruce con las propiedades del bracket de Kauffman.

Para una explicación más detallada del truco de Dirac, puede consultarse [34] y también el texto principal de Kauffman [18].

CONCLUSIONES

1. Un estudio riguroso del modelo sobre el cual está basada la teoría de nudos demanda una comprensión profunda de topología y más específicamente, motiva el estudio de la topología de variedades de dimensiones bajas, topología de variedades lineales a trozos y algunos aspectos de la topología de variedades orientables. A pesar de esto, un estudio exploratorio de la teoría de nudos puede ser realizado sin ningún conocimiento previo.
2. El grupo fundamental y conceptos como la equivalencia homotópica dan otra noción de equivalencia topológica, la cual puede ser explicada en términos intuitivos y dan lugar a conceptos profundos dentro de la topología algebraica, como el producto libre amalgamado y el teorema de Seifert-Van Kampen. En particular, el uso de este último teorema fue la parte más importante de la justificación para la presentación de Wirtinger usada para el cálculo de grupos de nudos y enlaces.
3. Se presentaron estructuras algebraicas (ya conocidas) asociadas a nudos y enlaces dentro de las categorías **Grp** con el grupo de un nudo, **Ring** con corchetes de Kauffman y **Ab** con la homología de Khovanov. Esto indica una amplia diversidad de estructuras algebraicas asociadas a nudos.
4. La exposición en este documento de las conexiones entre la física y la teoría de nudos, indica una gran generalidad para las construcciones presentadas, pues son aplicables en ambos contextos.

RECOMENDACIONES

1. Tomar como meta en un primer curso de topología general, un breve estudio de la teoría de nudos. Enfocar los esfuerzos del estudiante por comprender los conceptos expuestos en parte del capítulo 1, el capítulo 2 y el apéndice A, resulta más productivo que, simplemente, exponer las construcciones de manera abstracta.
2. Explorar las propiedades de la homología de Khovanov y su relación con el concepto general de homología.
3. Hacer referencia a la teoría de nudos en un estudio orientado a física de partículas o física de estado sólido en el que se traten los conceptos expuestos en el último capítulo.

BIBLIOGRAFÍA

- [1] ADAMS, COLIN CONRAD: *The knot book: an elementary introduction to the mathematical theory of knots*. W.H. Freeman, 1994.
- [2] AHLFORS, LARS: *Complex Analysis*. McGraw-Hill ScienceEngineeringMath, 3ª edición, 1979. ISBN 0070006571, 9780070006577.
- [3] BAXTER, R.J.: *Exactly Solved Models in Statistical Mechanics*. Academic Press, 1982.
- [4] BREDON, GLEN E: *Topology and geometry*. volumen 139. Springer Science & Business Media, 2013.
- [5] BROWN, RONALD: *Topology and Groupoids*. BookSurge, Scotts Valley, CA, 2007.
- [6] BURTON, BENJAMIN A.: «The Next 350 Million Knots». En: Sergio Cabello y Danny Z. Chen (Eds.), *36th International Symposium on Computational Geometry (SoCG 2020)*, volumen 164 de *Leibniz International Proceedings in Informatics (LIPIcs)*, pp. 25:1–25:17. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl, Germany. ISBN 978-3-95977-143-6. ISSN 1868-8969, 2020. doi: 10.4230/LIPIcs.SoCG.2020.25.
<https://drops.dagstuhl.de/entities/document/10.4230/LIPIcs.SoCG.2020.25>
- [7] CAIRNS, STEWART S.: «An Elementary Proof of the Jordan-Schoenflies Theorem». *Proceedings of the American Mathematical Society*, 1951, **2(6)**, pp. 860–867. ISSN 00029939, 10886826.
<http://www.jstor.org/stable/2031698>
- [8] CARTER, J. SCOTT y SAITO, MASAHIKO: «Reidemeister moves for surfaces isotopies and their interpretation as moves to movies». *Journal of Knot Theory and Its Ramifications*, 1993, **02(03)**, pp. 251–284. doi: 10.1142/S0218216593000167.

- [9] CARTER, J SCOTT y SAITO, MASAHIKO: *Knotted surfaces and their diagrams*. volumen 55. American Mathematical Society, 2023.
- [10] COLIN ROURKE, BRIAN SANDERSON: *Introduction to Piecewise-Linear Topology*. Springer Study Edition. Springer-Verlag Berlin Heidelberg, 11 edición, 1982. ISBN 9783540111023,9783642817359.
- [11] DROBOTUKHINA, YULIA V: «An analogue of the Jones polynomial for links in $\mathbb{R}P^3$ and a generalization of the Kauffman–Murasugi theorem». *Algebra i analiz*, 1990, **2(3)**, pp. 171–191.
- [12] FULTON, WILLIAM: *Algebraic topology: a first course*. Springer, 1995. ISBN 0387943277,9780387943275.
- [13] HATCHER, ALLEN: *Algebraic Topology*. Cambridge University Press, 2002.
- [14] HUDSON, J. F. P.: *Piecewise Linear Topology*. W. A. Benjamin, 1969.
- [15] KAUFFMAN, LOUIS H.: «Statistical mechanics and the Jones polynomial». *New problems, methods and techniques in quantum field theory and statistical mechanics*, 1988, pp. 175–222.
- [16] —: «A Tutte polynomial for signed graphs». *Discrete Applied Mathematics*, 1989, **25(1-2)**, pp. 105–127.
- [17] —: *The Interface of Knots and Physics: American Mathematical Society Short Course January 2-3, 1995 San Francisco, California (Proceedings of Symposia in Applied Mathematics)*. Amer Mathematical Society, 1995.
- [18] —: *Knots and physics*. World Scientific, 3rd edl edición, 2001. ISBN 978-058-54-5917-2,978-981-02-4111-7,978-981-02-4112-4.
<https://doi.org/10.1142/4256>
- [19] LEE, JOHN: *Introduction to topological manifolds*. volumen 202. Springer Science & Business Media, 2010.
- [20] LICKORISH, WB RAYMOND: *An introduction to knot theory*. volumen 175. Springer Science & Business Media, 1997.
- [21] LIMA, ELON LAGES: «Grupo fundamental e espaços de recobrimento, 2a edição». *IMPA, CNPq, Rio de Janeiro*, 1998.

- [22] —: *Homologia básica*. IMPA, 2009.
- [23] LIN, SIYANG: «Knots and Feynman Diagrams», 2018.
- [24] LIVINGSTON, CHARLES: «Connected sums of codimension two locally flat submanifolds», 2021.
<https://arxiv.org/abs/2104.03930>
- [25] MONTROYA-VEGA, GABRIEL: «Una Mirada Inicial a la Teoría de Nudos y a la Homología de Khovanov», 2023.
<https://arxiv.org/abs/2308.10277>
- [26] MUNKRES, JAMES: *Topology*. Prentice Hall, Inc, 2nd ed^{ta} edición, 2000. ISBN 9780131816299,0131816292.
- [27] NEGAMI, SEIYA: «Ramsey theorems for knots, links and spatial graphs». *Transactions of the American Mathematical Society*, 1991, **324**(2), pp. 527–541.
- [28] R. H. CROWELL, R. H. FOX: *Introduction to knot theory*. Graduate Texts in Mathematics. Springer, 1963 corr 4th printing^{ta} edición, 1984. ISBN 978-038-79-0272-2,978-354-09-0272-0.
<https://doi.org/10.1007/978-1-4612-9935-6>
- [29] REIDEMEISTER, KURT: «Elementare begründung der knotentheorie». En: *Abhandlungen aus dem Mathematischen Seminar der Universität Hamburg*, volumen 5, pp. 24–32. Springer, 1927.
- [30] ROGER A. HORN, CHARLES R. JOHNSON: *Matrix Analysis*. Cambridge University Press, 1990. ISBN 0521386322, 9780521386326, 0521305861.
- [31] ROLFSEN, DALE: *Knots and links*. 346. American Mathematical Soc., 2003.
- [32] ROTMAN, JOSEPH J.: *An Introduction to Algebraic Topology*. Graduate Texts in Mathematics 119. Springer-Verlag New York, 1^a edición, 1988.
- [33] SIMON, STEVEN H: «Topological Quantum: Lecture Notes and Proto-Book». *Unpublished prototype.[online]* Available at: <http://www-thphys.physics.ox.ac.uk/people/SteveSimon>, 2020, **26**.

- [34] STALEY, MARK: «Understanding quaternions and the Dirac belt trick». *European Journal of Physics*, 2010, **31(3)**, p. 467478. ISSN 1361-6404. doi: 10.1088/0143-0807/31/3/004.
<http://dx.doi.org/10.1088/0143-0807/31/3/004>
- [35] STILLWELL, JOHN: *Classical topology and combinatorial group theory*. Springer, 2nd edición, 1993.
- [36] TAO, TERENCE: «van Kampen's theorem via covering spaces», 2012.
<https://terrytao.wordpress.com/2012/10/28/van-kampens-theorem-via-covering-spaces/>
- [37] TEMPERLEY, H. N. V. y LIEB, E. H.: «Relations between the 'Percolation' and 'Colouring' Problem and other Graph-Theoretical Problems Associated with Regular Planar Lattices: Some Exact Results for the 'Percolation' Problem». *Proceedings of the Royal Society of London. Series A, Mathematical and Physical Sciences*, 1971, **322(1549)**, pp. 251–280.
<http://www.jstor.org/stable/77727>
- [38] VIRO, OLEG: «Remarks on definition of Khovanov homology», 2002.
<https://arxiv.org/abs/math/0202199>
- [39] —: «Khovanov homology, its definitions and ramifications». *Fundamenta Mathematicae*, 2004, **184(1)**, pp. 317–342.
<http://eudml.org/doc/282696>

A. TOPOLOGÍA GENERAL

Resulta muy productivo encaminar un primer acercamiento a la topología general a temas útiles para otras ramas de la matemática, algunos de los enfoques pueden ser, un acercamiento algebraico/categorico, otro orientado a la geometría diferencial de variedades suaves o un enfoque orientado a espacios de funciones. En el caso de este apéndice, se presentarán generalidades muy superficiales de los primeros dos enfoques. Una de las motivaciones para estudiar variedades es trabajar con espacios topológicos con una estructura similar a $\mathbb{R}^n$, por ello, necesitaremos algunos preliminares de topología general antes de definir que es una variedad.

A.1 Definición. Una *topología* en un conjunto X es una colección de subconjuntos $\mathcal{T}$ con las siguientes propiedades:

1. $\mathcal{T}$ tiene al menos al conjunto total y el conjunto vacío, más concretamente $X, \emptyset \in \mathcal{T}$.
2. $\mathcal{T}$ es cerrada bajo una cantidad arbitraria de uniones, es decir, si $\{U_\alpha\}_{\alpha \in A} \subseteq \mathcal{T}$ para algún conjunto de índices A entonces

$$\bigcup_{\alpha \in A} U_\alpha \in \mathcal{T}.$$

3. $\mathcal{T}$ es cerrada bajo una cantidad finita de intersecciones, esto es, si $U_1, \dots, U_n \in \mathcal{T}$ entonces

$$\bigcap_{i=1}^n U_i \in \mathcal{T}$$

Al par $(X, \mathcal{T})$, de un conjunto X y una topología en el mismo conjunto lo llamaremos *espacio topológico*.

Cuando nos refiramos a un espacio topológico, usualmente, sólo se escribe el conjunto X y se especifica la topología o el contexto deja claro cual es la colección $\mathcal{T}$.

A los elementos de $\mathcal{T}$ los llamaremos *conjuntos abiertos* y si en un espacio topológico tenemos $p \in X$ a un abierto U tal que $p \in U$ le llamaremos *vecindad de p* . Asimismo, un concepto importante en topología es el de función continua.

A.2 Definición. Si X e Y son dos espacios topológicos, diremos que $f : X \rightarrow Y$ es una *función continua* si para todo abierto $U \subseteq Y$ tenemos que $f^{-1}(U) \subseteq X$ también es abierto, en ambos casos, abiertos en sus respectivas topologías.

De manera similar que en espacios métricos, podemos definir conjuntos ampliamente relacionados con los abiertos.

A.3 Definición. Dado un espacio topológico X , un conjunto $V \subseteq X$ se dice *cerrado* si su complemento respecto al espacio topológico $V^c = X \setminus V$ es abierto.

A.4 Definición. Si $(X, \mathcal{T})$ es un espacio topológico y $A \subseteq X$ tendremos, que si

$$\mathcal{F} = \{F : F \text{ es cerrado en } \mathcal{T} \text{ y } A \subseteq F\}$$

se define la *clausura de A* como

$$\overline{A} = \bigcap_{F \in \mathcal{F}} F.$$

A los elementos de $\overline{A}$ se les conoce como *puntos de adherencia* o *puntos adherentes*.

Notemos que esta definición siempre será relativa a la topología en la que estemos trabajando, por eso algunas veces se opta por la notación $\text{cl}_{(X, \mathcal{T})}(A)$ para la clausura, sin embargo, a partir de ahora siempre quedará claro tanto la topología y el conjunto X sobre el cual estaremos trabajando por lo que se opta por la notación $\overline{A}$. En un primer curso de análisis, y en general cuando estamos trabajando en espacios métricos, usualmente, se define la clausura como

$$\overline{A} = A \cup \left\{ \lim_{n \rightarrow \infty} a_n : a_n \in A \text{ para todo } n \in \mathbb{N} \right\},$$

donde los límites y la definición de convergencia viene dada por la métrica subyacente del espacio. Otra definición bastante popular es la de que $\overline{A} = A \cup A'$ donde A' se le conoce como *el conjunto derivado* y consiste en todos los puntos límite de el conjunto A . Todas estas definiciones son equivalentes en espacios métricos.

A.5 Definición. Si $A \subseteq X$ se dice que un conjunto A es *denso en X* , si

$$\overline{A} = X.$$

Con esto definido, una caracterización bastante útil para funciones continuas es la siguiente:

A.6 Proposición. *Dados dos espacios topológicos X, Y se dice que $f : X \rightarrow Y$ es continua si y sólo si*

$$f(\overline{A}) \subseteq \overline{f(A)}$$

para todo $A \subseteq X$.

Un caso particular de los mapeos continuos son los considerados como los isomorfismos en topología, definidos a continuación.

A.7 Definición. Si X e Y son dos espacios topológicos, un *homeomorfismo* de X a Y es una función continua y biyectiva $\varphi : X \rightarrow Y$ cuya inversa también es continua. Si existe un homeomorfismo entre dos espacios topológicos diremos que estos son *homeomorfos* o *topológicamente equivalentes*.

Ahora, algo complicado es dar de manera explícita todos los conjuntos que pertenecen a una topología, por ello recurriremos a las bases, las cuales nos permiten dar una descripción detallada de nuestro espacio topológico.

A.8 Definición. Si X es un conjunto, una *base para la topología en X* es una colección $\mathcal{B}$ de subconjuntos de X (usualmente, a dichos subconjuntos se les llama básicos o elementos base) tal que:

1. Para cada $x \in X$, existe al menos un $B \in \mathcal{B}$ tal que $x \in B$.
2. Si $x \in B_1 \cap B_2$ para dos elementos base B_1 y B_2 , debe existir un tercer elemento base B_3 que verifique que $x \in B_3$ y además $B_3 \subseteq B_1 \cap B_2$.

Dado un conjunto X , con esta definición, si $\mathcal{B}$ satisface las dos condiciones para ser una base, definiremos a la *topología $\mathcal{T}$ generada por $\mathcal{B}$* de la siguiente forma: un subconjunto U de X se dice abierto en X o dicho de otra manera se tiene $U \in \mathcal{T}$ si y sólo si para cada $x \in U$ existe un elemento base B tal que $x \in B$ y $B \subseteq U$. Y como bien hemos nombrado a esta colección, esta es una topología y eso debemos demostrarlo.

A.9 Teorema. *Dado un conjunto X y una base $\mathcal{B}$ para una topología en el mismo conjunto, entonces $\mathcal{T}$, la topología generada por $\mathcal{B}$ es en efecto una topología.*

Si consideramos al conjunto vacío, tendremos que $\emptyset \in \mathcal{T}$ por vacuidad, si $x \in X$ entonces, por definición de base se tiene $X \in \mathcal{T}$ eso verifica la primera propiedad. Para la segunda propiedad tomemos una colección conjuntos $\{U_\alpha\}_{\alpha \in A}$ contenidos en la topología entonces si consideramos el conjunto

$$U = \bigcup_{\alpha \in A} U_\alpha$$

veamos que si $x \in U$ entonces $x \in U_\alpha$ para algún α y por tanto existe conjunto base B tal que $x \in B$ con $B \subseteq U_\alpha$, implicando esto que $B \subseteq U$. Finalmente, notemos que basta con probar que si $U_1, U_2 \in \mathcal{T}$, entonces $U_1 \cap U_2 \in \mathcal{T}$, veamos entonces que si $x \in U_1 \cap U_2$, entonces $x \in U_1$ y $x \in U_2$, por lo que $x \in B_1$ y $x \in B_2$ para básicos B_1, B_2 y por tanto $x \in B_1 \cap B_2$, de esto que debe existir un $B_3 \in \mathcal{B}$ tal que $x \in B_3 \subseteq B_1 \cap B_2 \subseteq U_1 \cap U_2$, obteniendo $U_1 \cap U_2 \in \mathcal{T}$. ■

Ahora presentaremos un resultado que deja un poco más clara la relación entre una topología y su base.

A.10 Teorema. Sea X un conjunto; $\mathcal{B}$ una base para la topología $\mathcal{T}$ en X . Entonces $\mathcal{T}$ es igual a la colección de todas las uniones de los elementos de $\mathcal{B}$.

Con estos conceptos, ya nos vamos acercarnos a lo que es una variedad. Ahora veamos una propiedad que nos interesa, pues hace que un espacio topológico se parezca localmente a $\mathbb{R}^n$.

A.11 Definición. Sea M un espacio topológico, este se dice *localmente euclidiano de dimensión n* si para todo $q \in M$ existe una vecindad de este punto que es homeomorfo a un subconjunto abierto de $\mathbb{R}^n$. Dicha vecindad es llamada vecindad euclídea de q .

La siguiente propiedad está motivada en que necesitamos suficientes abiertos en el espacio topológico en el cual estamos trabajando.

A.12 Definición. Un espacio topológico X es llamado un *espacio de Hausdorff* si separa puntos, esto es, si para cualesquiera dos puntos distintos $q_1, q_2 \in X$ existen dos vecindades de estos U_1, U_2 tales que $q_1 \in U_1$ y $q_2 \in U_2$, y además $U_1 \cap U_2 = \emptyset$.

Y finalmente, la siguiente definición viene de que requerimos suficientes abiertos en la topología, pero no demasiados.

A.13 Definición. Un espacio topológico se dice *segundo contable* si admite una base numerable.

Con esto, podemos definir lo que es una variedad topológica.

A.14 Definición. Una *variedad topológica* de dimensión n es un espacio topológico que es Hausdorff, segundo contable y localmente euclidiano de dimensión n .

Para abreviar, usaremos el término n -variedad cuando hablemos de una variedad topológica de dimensión n .

A continuación presentaremos algunas topologías que pueden ser definidas en función de otros espacios topológicos. El primero de ellos es la topología del subespacio.

A.15 Definición. Sea X un espacio topológico, y sea $S \subseteq X$ cualquier subconjunto. Definimos la colección de subconjuntos de S como:

$$\mathcal{T}_S = \{U \subseteq S : U = S \cap V \text{ para algún subconjunto } V \subseteq X\}.$$

A la colección $\mathcal{T}_S$ se le conoce como *topología del subespacio* o *topología relativa a S* .

El que esta colección sea en efecto una topología y $(S, \mathcal{T}_S)$ sea un espacio topológico es consecuencia directa de las propiedades de la intersección y unión de conjuntos. En este caso se puede mostrar que los cerrados de esta topología son los conjuntos de la forma $S \cap W$ con W cerrado en X . Con ello se tiene una propiedad característica para este espacio.

A.16 Proposición (Propiedad Característica de la Topología de Subespacio). *Sea X un espacio topológico y $S \subseteq X$ un subespacio. Para cualquier espacio topológico Y , un mapeo $f : Y \rightarrow S$ es continuo si y sólo si la composición $\iota_S \circ f : Y \rightarrow X$ es continua.*

$$\begin{array}{ccc} & & X \\ & \nearrow^{\iota_S \circ f} & \uparrow \iota_S \\ Y & \xrightarrow{f} & S. \end{array}$$

Con un poco más claridad en la topología de un subconjunto de otro espacio topológico, definimos un tipo especial de función continua.

A.17 Definición. Una función continua $\varphi : X \rightarrow Y$ es un *encaje topológico* si $\varphi : X \rightarrow \varphi(X)$ es un homeomorfismo.

También, un resultado que nos asegura que una función definida a trozos es continua es el siguiente:

A.18 Teorema (Lema del Pegado). *Sean X, Y espacios topológicos, y sea la familia indizada de conjuntos $\{A_\alpha : \alpha \in \Lambda\}$ tales que $X = \bigcup_{\alpha \in \Lambda} A_\alpha$. Si para cada $\alpha \in \Lambda$, $f_\alpha : A_\alpha \rightarrow Y$ son aplicaciones continuas (A_α con la topología de subespacio) tales que $f_\gamma|_{A_\gamma \cap A_\beta} = f_\beta|_{A_\gamma \cap A_\beta}$ para cada $\gamma, \beta \in \Lambda$. Si se considera la aplicación $f : X \rightarrow Y$ definida por $f|_{A_\alpha} = f_\alpha$, entonces f es continua si se cumple alguna de las siguientes condiciones:*

- a. *Todos los A_α son abiertos.*
- b. *Todos los A_α son cerrados y Λ es finito.*

Otra topología bastante utilizada en la teoría de variedades topológicas es la topología cociente.

A.19 Definición. Dado X un espacio topológico, Y un conjunto arbitrario, y una aplicación sobreyectiva $q : X \rightarrow Y$. Se considera la topología en Y definiendo a $U \subseteq Y$ como un conjunto abierto si y sólo si $q^{-1}(U)$ es abierto en X . A esta topología se le conoce como *topología cociente*.

A q en la definición la llamaremos *aplicación cociente*. Además esta topología también posee una propiedad que la caracteriza. Pues la idea de la topología cociente es el espacio más pequeño para el cual la aplicación q es continua.

A.20 Proposición (Propiedad Característica de la Topología Cociente). *Sean X, Y espacios topológicos y $q : X \rightarrow Y$ es un aplicación cociente. Para todo espacio topológico Z , una aplicación $f : Y \rightarrow Z$ es continua si y sólo si la composición $f \circ q$ es continua.*

$$\begin{array}{ccc} X & & \\ q \downarrow & \searrow f \circ q & \\ Y & \xrightarrow{f} & Z \end{array}$$

Esta propiedad característica puede ser usada (además de para demostrar la unicidad de la topología cociente salvo homeomorfismos) para mostrar el siguiente resultado, el cual es sumamente importante, pues nos da una manera de construir funciones continuas en espacios cociente, ya que son espacios en los que su estructura depende de otro espacio.

A.21 Teorema. Sea $q : X \rightarrow Y$ un mapeo cociente Z un espacio topológico, y $f : X \rightarrow Z$ cualquier función continua que es constante en las fibras de q (es decir, si $q(x) = q(x')$, entonces $f(x) = f(x')$). Entonces, existe una única aplicación continua $\tilde{f} : Y \rightarrow Z$ tal que el siguiente diagrama conmuta:

$$\begin{array}{ccc} X & & \\ q \downarrow & \searrow f & \\ Y & \xrightarrow{\tilde{f}} & Z \end{array}$$

Usualmente, a esta manera de definir funciones continuas se le conoce como *pasando el cociente*.

Otros espacios de suma importancia dentro de la topología son los compactos. Recordando primero, una *cubierta abierta* para un espacio X es una colección de conjuntos abiertos $\mathcal{U}$ tales que $X = \bigcup_{U \in \mathcal{U}} U$ y una *subcubierta* es una subcolección de $\mathcal{U}$, a la cual se denota por $\mathcal{V}$ tal que $X = \bigcup_{V \in \mathcal{V}} V$, un espacio X se dice *compacto* si toda cubierta abierta de X tiene una subcubierta finita, un subconjunto de un espacio topológico es llamado compacto si es un espacio compacto con la topología del subespacio. Una función se dice *aplicación cerrada* si mapea cerrados en cerrados. En concreto una de las proposiciones más importantes cuando se trabaja con compactos es la siguiente:

A.22 Proposición. Sea $F : X \rightarrow Y$ una función continua donde X es compacto e Y es espacio de Hausdorff, entonces:

- F es una aplicación cerrada.
- Si F sobreyectiva, F es una aplicación cociente.
- Si F inyectiva, F es un encaje topológico.
- Si F biyectiva, F es un homeomorfismo.

Notemos que una relación de equivalencia también puede inducir una topología cociente.

A.23 Definición. Si $\sim$ es una relación de equivalencia en un espacio topológico X , al conjunto $X/\sim$ junto con la topología cociente inducida por la proyección canónica $q : X \rightarrow X/\sim$, se le conoce como *espacio de identificación* o también *espacio cociente*.

Recordando que una relación de equivalencia también puede ser dada por una partición de un conjunto, definimos el siguiente espacio.

A.24 Definición. Si $X_1, \dots, X_k$ son espacios topológicos no vacíos. Para cada i , sea p_i un punto específico en X_i ; se define la relación de equivalencia $\sim$ inducida por la partición de $\bigsqcup_{i=1}^k X_i$ dada por:

$$\{(p_1, 1), \dots, (p_k, k)\} \cup \left\{ \{(p, i)\} : (p, i) \in \bigsqcup_{i=1}^k X_i, (p, i) \neq (p_j, j), \text{ para } j = 1, \dots, k \right\},$$

al espacio de identificación

$$\bigvee_{i=1}^k X_i = X_1 \vee \dots \vee X_k = \left(\bigsqcup_{i=1}^k X_i \right) / \sim,$$

se le conoce como *producto cuña* de los espacios $X_1, \dots, X_k$.

El lector podrá inferir una definición análoga para $\bigvee_{\alpha} X_{\alpha}$ cuando α varíe en un conjunto arbitrario de índices. Además, en la mayoría de espacios esta elección de los puntos p_i , los cuales son llamados *puntos base*, no es relevante, por lo que esta notación hace sentido, en el caso en el que la elección de puntos base sea relevante se hace la aclaración. En particular, estos espacios son de suma importancia para la construcción de espacios con propiedades algebraicas relevantes.

B. HOMOLOGÍA

A manera de estudio complementario, se presentaran algunas nociones generales de homología, enteramente basadas en el libro [22] de Lima, originalmente publicado en portugués. En esta última referencia se presentan algunos casos particulares de la teoría de homología, como la homología singular o la homología simplicial.

B.1. Definiciones generales

Antes de entrar de lleno con las definiciones generales en homología revisaremos una construcción particular que también aparece en la homología de Khovanov. Existe una manera de construir grupos mediante conjuntos arbitrarios, estos están ampliamente relacionados con el concepto de grupos libres y serán muy útiles en este contexto.

B.1 Definición. Sea $\mathcal{S}$ un conjunto cualquiera, se define $\mathbb{Z}\langle\mathcal{S}\rangle$ como el conjunto de todas las funciones $\varphi : \mathcal{S} \rightarrow \mathbb{Z}$ tales que $\varphi(s) = 0$ para todos salvo una cantidad finita de elementos $s \in \mathcal{S}$. Así, $\mathbb{Z}\langle\mathcal{S}\rangle$ es el *grupo abeliano libre* con base $\mathcal{S}$.

Usualmente, se escribe como $k \cdot x$ o kx a la función $\varphi \in \mathbb{Z}\langle\mathcal{S}\rangle$ tal que $\varphi(x) = k$ y $\varphi(y) = 0$ para todo $y \neq x$. Además, puede mostrarse que cada φ tiene una única representación de la forma:

$$\varphi = \sum_{x \in \mathcal{S}} k_x \cdot x.$$

De esta manera se identificará a $x \in \mathcal{S}$ en $\mathbb{Z}\langle\mathcal{S}\rangle$ con $\mathbf{1}_{\{x\}} = 1 \cdot x$. Otra notación usada para $\mathbb{Z}\langle\mathcal{S}\rangle$ es $F^{\text{ab}}(\mathcal{S})$ pues este grupo puede ser visto como la abelianización del grupo $F(\mathcal{S})$, el grupo libre en $\mathcal{S}$ descrito en el capítulo 2. Recordando que la

abelianización de un grupo G se define como el cociente:

$$G^{\text{ab}} = \frac{G}{[G, G]},$$

del grupo y su grupo de conmutadores.

En general, un complejo de cadena puede ser definido mediante un cierto tipo de anillos bastante común en álgebra. En el caso de la homología de Khovanov, un grupo abeliano libre puede verse como un $\mathbb{Z}$ -módulo, pero en general podemos cambiar $\mathbb{Z}$ por otra estructura.

B.2 Definición. Sea A un anillo conmutativo con unidad. Un *complejo de cadena* con coeficientes en A es una secuencia $\mathcal{C} = \{(\mathcal{C}_p, \partial_p)\}$ de A -módulos $\mathcal{C}_p$, $p \geq 0$ entero y homomorfismos $\partial_p : \mathcal{C}_p \rightarrow \mathcal{C}_{p-1}$ tales que $\partial_p \circ \partial_{p+1} = 0$. Esto se escribe:

$$\mathcal{C} : \cdots \rightarrow \mathcal{C}_{p+1} \xrightarrow{\partial_{p+1}} \mathcal{C}_p \xrightarrow{\partial_p} \mathcal{C}_{p-1} \rightarrow \cdots \rightarrow \mathcal{C}_1 \xrightarrow{\partial_1} \mathcal{C}_0 \xrightarrow{\partial_0} 0.$$

Cada elemento $x \in \mathcal{C}_p$ es llamado *p-cadena*, o *cadena de dimensión p* y cada homomorfismo ∂_p es llamado *operador frontera*.

Si $\partial_p x = 0$ se dice que x es un *p-ciclo* o simplemente ciclo. El conjunto de los *p-ciclos* se denota, usualmente, por Z_p y es un submódulo de $\mathcal{C}_p$, pues $Z_p = \text{Ker } \partial_p$. Si $y = \partial_{p+1}x$, se dice que la *p-cadena y* es la frontera de la $(p+1)$ -cadena x . El conjunto B_p de las *p-cadenas* que son fronteras de alguna $(p+1)$ -cadena, es también un submódulo de $\mathcal{C}_p$ pues $B_p = \text{im } \partial_{p+1}$. Muchas veces se deja implícito y se opta por la notación ∂ en lugar de ∂_p (esto pasa por ejemplo cuando se trabaja con formas diferenciales). La relación $\partial_p \circ \partial_{p+1} = 0$ de la Definición B.2 implica que $B_p \subset Z_p$, lo cual justifica la siguiente definición:

B.3 Definición. Con la notación empleada hasta ahora, el A -módulo cociente

$$H_p(\mathcal{C}) := \frac{Z_p}{B_p},$$

es llamado *grupo de homología p-dimensional* del complejo $\mathcal{C}$ con coeficientes en A o la *homología* del complejo $\mathcal{C}$. A los elementos del cociente

$$[z] = z + B_p,$$

se les conoce como *clases de homología*.

B.4 Definición. Sean $\mathfrak{X} = \{(X_p, \partial_p^X)\}$ y $\mathfrak{Y} = \{(Y_p, \partial_p^Y)\}$ complejos de cadena. Un morfismo $f : \mathfrak{X} \rightarrow \mathfrak{Y}$ es una secuencia de homomorfismos $f_p : X_p \rightarrow Y_p$ tales que $f_p(\partial_p^X x) = \partial_p^Y f_{p+1}(x)$ para todo $x \in X_{p+1}$. En otras palabras, todos los rectángulos del siguiente diagrama conmutan:

$$\begin{array}{ccccccccccc} \cdots & \longrightarrow & X_{p+1} & \xrightarrow{\partial_{p+1}^X} & X_p & \xrightarrow{\partial_p^X} & X_{p-1} & \xrightarrow{\partial_{p-1}^X} & \cdots & \xrightarrow{\partial_1^X} & X_0 \\ & & \downarrow f_{p+1} & & \downarrow f_p & & \downarrow f_{p-1} & & & & \downarrow f_0 \\ \cdots & \longrightarrow & Y_{p+1} & \xrightarrow{\partial_{p+1}^Y} & Y_p & \xrightarrow{\partial_p^Y} & Y_{p-1} & \xrightarrow{\partial_{p-1}^Y} & \cdots & \xrightarrow{\partial_1^Y} & Y_0. \end{array}$$

Notemos que para diferenciar los operadores frontera en cada complejo de cadena se usa la notación $\partial_p^X, \partial_p^Y$, sin embargo esto puede ser omitido y escribir ∂ en ambos casos, pues queda implícito el dominio de cada operador frontera. En resumen, para dos complejos de cadenas tendremos:

$$\begin{array}{ccccccccccc} \cdots & \longrightarrow & X_{p+1} & \xrightarrow{\partial} & X_p & \xrightarrow{\partial} & X_{p-1} & \xrightarrow{\partial} & \cdots & \xrightarrow{\partial} & X_0 \\ & & \downarrow f_{p+1} \curvearrowright & & \downarrow f_p \curvearrowright & & \downarrow f_{p-1} \curvearrowright & & & & \downarrow f_0 \curvearrowright \\ \cdots & \longrightarrow & Y_{p+1} & \xrightarrow{\partial} & Y_p & \xrightarrow{\partial} & Y_{p-1} & \xrightarrow{\partial} & \cdots & \xrightarrow{\partial} & Y_0. \end{array}$$

B.2. Homología singular

Ahora presentaremos uno de los tipos más intuitivos de homología, el cual es usado para generalizar el concepto de conexidad en espacios topológicos.

B.5 Definición. Sea X un espacio topológico, un n -simplejo singular en X es una aplicación continua $\sigma : \Delta^n \rightarrow X$ del n -simplejo estándar al espacio topológico.

Definimos nuestros grupos de n -cadenas y operadores frontera para la construcción de los grupos de homología.

B.6 Definición. Sea X un espacio topológico. Para cada $n \geq 0$, se define $S_n(X)$ como el grupo abeliano libre con todos los n -simplejos singulares en X como base; así también $S_{-1}(X)$ se toma por definición como el grupo trivial. Los elementos de $S_n(X)$ son llamados n -cadenas singulares en X .

B.7 Definición. Si $\sigma : \Delta^n \rightarrow X$ continua y $n > 0$, se define su frontera por

$$\partial_n \sigma = \sum_{i=0}^n (-1)^i \sigma \circ \varepsilon_i^n = \sum_{i=0}^n (-1)^i \sigma \varepsilon_i^n \in S_{n-1}(X),$$

donde $\varepsilon_i^n : \Delta^{n-1} \rightarrow \Delta^n$ se define como

$$\begin{aligned}\varepsilon_0^n(t_1, \dots, t_{n-1}) &= (0, t_1, \dots, t_{n-1}), \\ \varepsilon_i^n(t_1, \dots, t_{n-1}) &= (t_1, \dots, t_{i-1}, 0, t_i, \dots, t_{n-1}),\end{aligned}$$

y para $n = 0$ se define $\partial_0 \sigma = 0$.

Podemos notar que ahora tenemos una sucesión $S_*(X)$ de $\mathbb{Z}$ -módulos y homomorfismos $\partial_n : S_n(X) \rightarrow S_{n-1}(X)$ dados por $\partial_n(\sigma) = \partial_n \sigma$ tales que

$$S_*(X) : \cdots \rightarrow S_{n+1} \xrightarrow{\partial_{n+1}} S_n \xrightarrow{\partial_n} S_{n-1} \rightarrow \cdots \rightarrow S_1 \xrightarrow{\partial_1} S_0 \xrightarrow{\partial_0} 0,$$

y $Z_n(X) = \text{Ker } \partial_n \subseteq \text{im } \partial_{n+1} = B_n(X)$, siendo esto último consecuencia directa de la definición de los operadores frontera. Dando pie a definir los grupos que nos interesan con el complejo de cadena $S_*(X) = \{(S_n(X), \partial_n)\}$.

B.8 Definición. Con la notación usada hasta ahora, si X es un espacio topológico, para cada $n \geq 0$, el n -ésimo grupo de homología $H_n(X)$ singular se define como:

$$H_n(X) = \frac{Z_n(X)}{B_n(X)}.$$

Siendo este un caso particular de las definiciones de la primera sección de este apéndice. Los morfismos en este caso pueden ser contruidos de la siguiente forma: si $f : X \rightarrow Y$ es una función continua entre espacios topológicos, para todo n -simplejo singular σ en X existe $f \circ \sigma$ es un n -simplejo singular en Y . Esta correspondencia induce homomorfismos $f_{\#,n} : S_n(X) \rightarrow S_n(Y)$ los cuales verifican $\partial_n f_{\#,n} = f_{\#,(n-1)} \partial_n$ y por ende, esta secuencia $f_{\#} : S_*(X) \rightarrow S_*(Y)$ de homomorfismos, es un morfismo. Finalmente, presentaremos una conexión sumamente interesante entre el grupo fundamental y el primer grupo de homología singular.

B.9 Teorema (Teorema de Hurewicz). *Sea X un espacio conexo por caminos y $p \in X$, entonces*

$$\pi_1(X, p)^{\text{ab}} \cong H_1(X),$$

donde $\pi_1(X, p)^{\text{ab}}$ es la abelianización del grupo fundamental de X con base en p .

Para los detalles sobre la prueba de este teorema y material introductorio a homología, el cual fue tomado como referencia para la construcción de este apéndice puede consultarse [32].