

UNIVERSIDAD AUTÓNOMA DE BAJA CALIFORNIA
Facultad de Ciencias



EL TEOREMA DE DEHN-NIELSEN-BAER
T E S I S
QUE PARA OBTENER EL TÍTULO DE
LICENCIADO EN MATEMÁTICAS APLICADAS
PRESENTA:

JOSÉ JOAQUÍN DOMÍNGUEZ SÁNCHEZ

DIRECTOR: DR. JESÚS HERNÁNDEZ HERNÁNDEZ

CODIRECTOR: DR. CARLOS YEE ROMERO

Ensenada, B.C.

Noviembre del 2020

UNIVERSIDAD AUTÓNOMA DE BAJA CALIFORNIA

Facultad de Ciencias

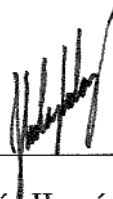
EL TEOREMA DE DEHN-NIELSEN-BAER

TESIS PROFESIONAL

QUE PRESENTA

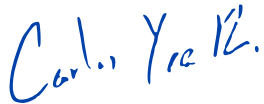
JOSÉ JOAQUÍN DOMÍNGUEZ SÁNCHEZ

APROBADO POR:



Dr. Jesús Hernández Hernández

Presidente del jurado



Dr. Carlos Yee Romero
SECRETARIO



Dra. Brenda Leticia De La Rosa Navarro
1ER. VOCAL

Agradecimientos

Primeramente, agradezco a mi asesor de tesis, el Dr. Jesús Hernández Hernández, por su apoyo, dedicación y paciencia al asesorarme en este trabajo. Sin él, este trabajo no hubiese sido posible.

Agradezco a la DGAPA-UNAM la beca recibida mediante el proyecto de Investigación realizada gracias al Programa de Apoyo a Proyectos de Investigación e Innovación Tecnológica (PAPIIT) de la UNAM IA104620.

También agradezco a mis sinodales, el Dr. Carlos Yee Romero y la Dra. Brenda Leticia De La Rosa Navarro, por sus aportaciones, sugerencias y correcciones para mejorar este trabajo de tesis.

Agradezco inmensamente a mi madre, padre, hermanas y hermano, por su apoyo, consejos, regaños, abrazos, amor y cariño. Un agradecimiento especial a mis cuñados.

A todos quienes han sido mis profesores, en especial, a los miembros de Academia de Matemáticas de Facultad de Ciencias de la UABC, por haberme enseñado y compartido sus conocimientos.

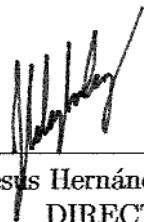
A mis compañeros y amigos de la universidad. No hubiera sido lo mismo sin ellos.

Resumen

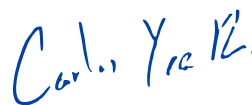
EL TEOREMA DE DEHN-NIELSEN-BAER

El grupo modular de Teichmüller fue introducido por primera vez en 1920 por el matemático alemán Max Dehn [10]. El *teorema de Dehn-Nielsen-Baer*, es un resultado de la interacción entre la topología y el álgebra en el grupo modular, cuya primera demostración conocida fue publicada por Jakob Nielsen en 1927 [2]. Tal teorema declara que el grupo modular de Teichmüller de una superficie cerrada orientable S_g , denotado por $\text{Mod}(S_g)$, es isomorfo a un subgrupo de índice dos del grupo automorfismos externos de $\pi_1(S_g, p)$, denotado por $\text{Out}(\pi_1(S_g))$. En este trabajo de tesis se presenta una demostración del teorema de Dehn-Nielsen-Baer para una superficie cerrada S_g de género $g \geq 2$ utilizando cuasi-isometrías. Para lograr esto se ha incurrido en el estudio de las propiedades topológicas y geométricas de una superficie, como por ejemplo, el grupo fundamental, espacios cubrientes, acciones de grupos, y el estudio de los modelos del plano hiperbólico, entre otros.

Resumen aprobado por:



Dr. Jesús Hernández Hernández
DIRECTOR



Dr. Carlos Yee Romero
CODIRECTOR

Índice general

Introducción	VII
1. Preliminares	1
1.1. Grupos	1
1.1.1. Acciones de grupos	2
1.1.2. Clases de Conjugación	5
1.1.3. Grupo de Automorfismos Externos	5
1.2. Superficies Hiperbólicas	6
1.2.1. Superficies	6
1.2.2. Superficies Hiperbólicas	10
1.3. Curvas Cerradas Simples	12
1.3.1. Curvas Cerradas	12
1.3.2. Curvas Cerradas Simples	14
1.3.3. Números de Intersección	19
1.3.4. Criterio del Bígono	20
1.4. Principio De Cambio de Coordenadas	22
1.4.1. Clasificación de curvas cerradas simples.	22
2. Grupo Modular de Teichmüller	27
2.1. Clases de Isotopía	27

2.2.	Grupo Modular de Teichmüller	30
2.3.	Método de Alexander	32
3.	Lema de Švarc-Milnor	36
3.1.	Cuasi-isometrías	37
3.1.1.	Encajes isométricos	37
3.1.2.	Encajes Cuasi-Isométricos	37
3.1.3.	Aplicaciones en Grupos Finitamente Generados	42
3.2.	Espacios Cuasi-Geodésicos	45
3.2.1.	Segmentos Geodésicos	45
3.2.2.	Segmentos Cuasi-Geodésicos	46
3.3.	Lema de Švarc-Milnor	48
4.	Ligados en el Infinito	57
4.1.	Métricas en $\pi_1(\mathcal{S}_g, p)$	57
4.2.	Ligados en el Infinito	59
5.	Teorema de Dehn-Nielsen-Baer	66
5.1.	Teorema de Dehn-Nielsen-Baer	66
5.1.1.	Grupo Modular de Teichmüller Extendido	66
5.1.2.	Grupo de Automorfismos Externos	67
5.1.3.	Homomorfismo Natural	68
5.1.4.	Teorema de Dehn-Nielsen-Baer	72
6.	Conclusión	80
A.	Geometría Hiperbólica	83
A.1.	Plano Hiperbólico	83
A.1.1.	Modelos de Plano Hiperbólico	83

A.1.2.	Transformaciones de Möbius	85
A.1.3.	Clasificación de Isometrías de $\mathbb{H}$	86
B.	Grupo Fundamental y Espacios Cubrientes	88
B.1.	Grupo Fundamental	88
B.1.1.	Caminos y Espacios Arco-conexos	89
B.1.2.	Homotopía de Funciones Continuas	89
B.1.3.	Producto de Caminos	91
B.1.4.	Grupo Fundamental	92
B.2.	Espacios Cubrientes	94
B.2.1.	Definición y Ejemplos de Espacios Cubrientes	95
B.2.2.	Levantamiento de Caminos	96
B.2.3.	Transformaciones De Cubierta	97
B.2.4.	La Acción de $\pi_1(X, x)$ en el Conjunto $\rho^{-1}(x)$	98

Introducción

Fue en el siglo III a.C. con el griego Euclides que se inició el estudio sistemático de la geometría plana en el tratado conocido como *Los Elementos*. Hasta ese momento la geometría era un conjunto de resultados, si bien formales, poco estructurados. Con Euclides, pasó a ser una ciencia deductiva; partiendo de ciertos conocimientos prefijados, y mediante un proceso lógico, se concluyen resultados [1].

Euclides parte de cinco axiomas, 23 definiciones y varias suposiciones implícitas, con los cuales deduce la mayor parte de los resultados sobre la geometría plana. El más característico de los postulados de Euclides es el equivalente al axioma de las paralelas: *Si una recta corta a otras dos de tal modo que los ángulos interiores del mismo lado suman menos que dos rectos, al prolongar indefinidamente estas dos rectas, se cortan en ese mismo lado.*

Durante 23 siglos los matemáticos creyeron que el quinto postulado se podría demostrar basándose en los otros cuatro, pero los esfuerzos para probarlo fueron inútiles. Fue en estos intentos de prueba que diversos matemáticos, entre ellos C.F. Gauss (1777 – 1855), J. Bolyai (1802 – 1860), N.I. Lobachevsky (1792 – 1856), E. Beltrami (1835 – 1899) y B. Riemann (1826 – 1866), introdujeron nuevas ideas que mostraron que era posible la existencia de otras geometrías distintas de la euclidiana. Con estas nuevas ideas se dio inicio a una nueva geometría no Euclidiana, conocida como geometría de Lobachevski o geometría hiperbólica; una geometría en donde las “paralelas” no son únicas [1].

Si bien ya se había aceptado la existencia de geometrías distintas a la euclidiana, no existía una manera estructurada que permitiera obtener resultados formales. Con el Programa de Erlangen, presentado por Felix Klein en 1872, se pudo clasificar y caracterizar a las geometrías existentes sobre la base de la geometría proyectiva y la teoría de grupos. Klein propuso que estudiar la geometría de un espacio consiste en estudiar las propiedades que permanecen invariantes bajo la acción de un grupo [1].

El programa de Erlangen, en muchos sentidos, inspiró a lo que hoy se conoce como la teoría geométrica de grupos, la cual es un área que consiste en el estudio de los grupos de dos maneras: mediante el estudio de sus acciones en objetos geométricos y a través del estudio de la geometría de los grupos como objetos geométricos en sí mismos [9]. La clave para unir estas dos formas de estudiar a los grupos es el Lema de Švarc Milnor, a través de la geometría a gran escala (cuasi-isometrías), burdamente este lema dice que dada una acción “bonita” (por isometrías, cocompacta y propiamente discontinua) de un grupo G sobre un espacio métrico “bonito” (propio y geodésico) X , se puede concluir que G es finitamente generado y cuasi-isométrico a X [9].

Si se quiere estudiar la geometría de una superficie S , eventualmente se llega al estudio del grupo fundamental de S , denotado por $\pi_1(S, p)$, y al estudio del grupo modular de Teichmüller de S , denotado por $\text{Mod}(S)$, el cual consiste de todos los homomorfismos de S que preservan la orientación módulo isotopías [4].

El resultado principal de esta tesis es el teorema de Dehn-Nielsen-Baer, el cual declara que, para una superficie S_g de género $g \geq 1$, el $\text{Mod}(S_g)$ es isomorfo a un subgrupo de índice 2 del $\text{Out}(\pi_1(S_g, p))$, el grupo de automorfismo externos del $\pi_1(S_g, p)$. Éste es un buen ejemplo de la interacción entre la topología y el álgebra en el grupo modular, pues relaciona un objeto de origen puramente topológico, $\text{Mod}(S_g)$, con un objeto de origen algebraico/topológico, $\text{Out}(\pi_1(S_g))$. Otras de las cosas más bellas de este resultado, es que para dar una primera demostración, además de las herramientas topológicas y algebraicas, también se utilizaron herramientas geométricas [2].

El objetivo de este trabajo es dar una demostración del teorema de Dehn-Nielsen-Baer utilizando cuasi-isometrías para cuando $g \geq 2$. Para esto, esta tesis se ha dividido en 6 capítulos y 2 apéndices.

En el Capítulo 1 primero se hace una recopilación de algunos resultados básicos de la teoría de grupos. Luego se establecen las bases para trabajar con curvas cerradas simples en una superficie. Se utiliza a la geometría hiperbólica para obtener que cada clase de homotopía de curvas cerradas simples en una superficie de género $g \geq 2$ tiene un representante geodésico único. Se introducen los conceptos de número de intersección geométrico y número intersección algebraico. Finalmente se presenta el principio de cambio de coordenadas, el cual desempeña el mismo papel que el de la matriz de cambio de base de un espacio vectorial.

El Capítulo 2 es una introducción al grupo modular de Teichmüller. En la primera parte se definen las clases de isotopías de curvas cerradas simples y las clases de isotopías de homeomorfismos. Finalmente se presenta el método de Alexander, el cual provee de un procedimiento para determinar si dos elementos de $\text{Mod}(S)$ son iguales. En particular, este método es usado para mostrar que un elemento de $\text{Mod}(S)$ es trivial o para verificar relaciones en $\text{Mod}(S)$. Este método se usa en la demostración del teorema de Dehn-Nielsen-Baer.

El Capítulo 3 inicia con algunas generalizaciones de isometrías y cuasi-isometrías. Se enuncian y demuestran dos versiones del Lema de Švarc Milnor, una versión métrica y una versión más topológica.

El Capítulo 4 es el último paso antes de comenzar con el teorema de Dehn-Nielsen-Baer. Se introduce la relación de ligaduras en infinito en el grupo fundamental y se prueba que esta relación es invariante bajo automorfismos. Este es un claro ejemplo del poder del Lema de Švarc Milnor, pues permite utilizar que cualquier automorfismo del $\pi_1(S_g, p)$ con la métrica hiperbólica es una cuasi-isometría.

En el Capítulo 5 se desarrolla la parte principal de este trabajo: el teorema de

Dehn-Nielsen-Baer. Se demuestra que el grupo modular de Teichmüller extendido de una superficie actúa de forma natural sobre el grupo fundamental mediante automorfismos externos. Se prueba que esta acción es un homomorfismo de grupos inyectivo y sobreyectivo, con lo que se demuestra el teorema de Dehn-Nielsen-Baer.

El Capítulo 6 es una recopilación de los resultados más relevantes para la demostración del teorema de Dehn-Nielsen-Baer. Se mencionan algunos problemas que surgen al estudiar al grupo Modular.

Para reforzar el contenido se han agregado dos apéndices. En el Apéndice A se recopilan resultados básicos de la geometría hiperbólica. El Apéndice B contiene la definición y algunas propiedades del grupo fundamental; una introducción corta de la teoría de espacios cubrientes; y la acción natural del grupo fundamental de un espacio, sobre su cubriente universal.

Capítulo 1

Preliminares

Este capítulo está destinado para realizar una recopilación de conceptos necesarios para un buen entendimiento de este trabajo. En la primera parte se presentan algunos resultados de la teoría de grupos. En la segunda (un poco más extensa) se incluyen resultados y ejemplos de superficies y superficies hiperbólicas. Seguidamente, se introduce el concepto curvas cerradas simples y se incluyen resultados relacionados a estas. Al final del capítulo se presenta el principio de cambio de coordenadas, una técnica frecuentemente utilizada en la teoría del grupo modular de Teichmüller (Capítulo 2).

1.1. Grupos

Recuérdese que un grupo G es un conjunto con una operación binaria $*$: $G \times G \rightarrow G$ que satisface tres propiedades: *asociatividad*, es decir, para todo $a, b, c \in G$ se cumple que $*(*(a, b), c) = *(a, *(b, c))$; *existencia (y unicidad) de elemento neutro*, es decir, existe (un único) $e \in G$ tal que $*(e, a) = a = *(a, e)$ para todo $a \in G$; y *existencia (y unicidad) de inversos*, es decir, para cada $a \in G$, existe (un único) $a^{-1} \in G$ tal que $*(a, a^{-1}) = e = *(a^{-1}, a)$. Habitualmente, a $*(a, b)$ se le denota

por $a * b$, o simplemente, ab si sobreentiende la operación. La asociatividad permite escribir a $(a * b) * c = a * (b * c)$ simplemente como $a * b * c$.

Definición 1.1. Sea G un grupo y $A \subset G$ un subconjunto. El **subgrupo generado por A (en G)**, denotado por $\langle A \rangle_G$ (o simplemente $\langle A \rangle$ si G se sobreentiende), es el mínimo subgrupo de G que contiene a A . Si $\langle A \rangle_G = G$ se dice que A es un **conjunto generador de G** , y si además A es finito, entonces G es **finitamente generado**.

Nótese que para todo grupo G y cualquier $A \subset G$, el subgrupo generado por A en G siempre existe y se puede describir como

$$\begin{aligned}\langle A \rangle_G &= \cap \{H \mid H \subset G \text{ es un subgrupo y } A \subset H\}, \\ &= \{a_1 * \cdots * a_n \mid n \in \mathbb{N}, a_1, \dots, a_n \in A \cup A^{-1}\},\end{aligned}$$

donde A^{-1} denota el conjunto que contiene a los inversos de los elementos de A .

1.1.1. Acciones de grupos

Definición 1.2. Una **acción de un grupo G en un conjunto X** es una función

$$\begin{aligned}\cdot : G \times X &\rightarrow X, \\ (g, x) &\mapsto g \cdot x\end{aligned}$$

tal que para todo $g, h \in G$ y todo $x \in X$ se cumple que $(g * h) \cdot x = g \cdot (h \cdot x)$ y $e \cdot x = x$, donde e es el elemento neutro de G .

Definición 1.3. Una **representación de un grupo G** es un homomorfismo de

grupos de G al grupo de biyecciones de un conjunto X , es decir, una función

$$\varphi : G \rightarrow \text{Biy}(X),$$

$$g \mapsto \varphi_g$$

tal que para todo $g, h \in G$, se satisface que $\varphi_{gh} = \varphi_g \circ \varphi_h$.

Nótese que si se define $g \cdot x := \varphi_g(x)$, entonces es equivalente tener una representación $\varphi : G \rightarrow \text{Biy}(X)$ y tener una acción de G en X . Si X es un espacio métrico y para cada $g \in G$, la función φ_g es una isometría, se dice que G **actúa por isometrías sobre X** . De manera análoga, si X es un espacio topológico y las biyecciones φ_g son todos homeomorfismos, se dice que G **actúa por homeomorfismos sobre X** .

Definición 1.4. Sea $\cdot : G \times X \rightarrow X$ una acción de G en X . La **órbita** de un elemento $x \in X$ respecto a esta acción es el conjunto

$$G \cdot x := \{g \cdot x \mid g \in G\}.$$

El **espacio de órbitas** de X (o el **cociente**) por la acción $\cdot$ es el conjunto

$$G \backslash X := \{G \cdot x \mid x \in X\}.$$

Definición 1.5. Sea G un grupo que actúa en un conjunto X . El **estabilizador** de un elemento $x \in X$ respecto a esta acción, es el subgrupo de G

$$G_x := \{g \in G \mid g \cdot x = x\}.$$

El **conjunto de puntos fijos** de un elemento $g \in G$ es el conjunto

$$X^g := \{x \in X \mid g \cdot x = x\}.$$

En particular, si $X^g = \emptyset$, se dice que g actúa sin puntos fijos.

Definición 1.6. Sean G un grupo y X un conjunto. Se dice que una acción de G en X es **libre** si para cada $g \in G$, con $g \neq e$, se tiene que

$$g \cdot x \neq x$$

para todo $x \in X$. En otras palabras, una acción es libre si y sólo si cada elemento no trivial del grupo actúa sin puntos fijos.

Definición 1.7. Una acción de G en X es **propriadamente discontinua** si es una acción por homeomorfismos tal que para todo subconjunto compacto $K \subset X$, el conjunto

$$\{g \in G \mid g \cdot K \cap K \neq \emptyset\}$$

es finito. En este caso, también se dice que G *actúa de forma propriadamente discontinua sobre X* . Aquí (y en lo consiguiente) $g \cdot K$ denota al conjunto $\{g \cdot k \mid k \in K\}$.

Definición 1.8. Una acción de G en X es **cocompacta** si es una acción por homeomorfismos para el cual existe un compacto $K \subset X$, tal que las G -traslaciones de K cubren a X , es decir,

$$G \cdot K := \bigcup_{g \in G} g \cdot K = X.$$

Definición 1.9. Sea X un espacio métrico, y sea G un grupo que actúa por homeomorfismos sobre X de forma propriadamente discontinua. Se dice que un subconjunto cerrado $E \subset X$ (es decir, la cerradura de un conjunto abierto $\overset{\circ}{E}$, llamado el interior de E), es un **dominio fundamental** para G si satisface

$$i) \bigcup_{g \in G} g \cdot E = X,$$

$$ii) \overset{\circ}{E} \cap g \cdot \overset{\circ}{E} = \emptyset \text{ para todo } g \in G - \{e\}.$$

1.1.2. Clases de Conjugación

Definición 1.10 (Conjugación). Sean a y b dos elementos de un grupo G . Se dice que b es **conjugado de a en G** , si existe $c \in G$ tal que $b = cac^{-1}$.

Nótese que si b es conjugado de a , entonces a es conjugado de b ; por tanto se puede decir simplemente que a y b son **conjugados** si alguna de estas dos relaciones se cumple. Más aún, esta relación es una relación de equivalencia. Las clases de equivalencia bajo esta relación se llaman **clases de conjugación**.

1.1.3. Grupo de Automorfismos Externos

Definición 1.11. Sea G un grupo. Un automorfismo $\varphi : G \rightarrow G$ es un **automorfismo interno de G** si φ está dado por una conjugación de un elemento de G , es decir, si existe un elemento $g \in G$ tal que

$$\varphi(h) = ghg^{-1}$$

para todo $h \in G$.

El **grupo de automorfismo internos de G** , denotado por $\text{Inn}(G)$, es el subgrupo del $\text{Aut}(G)$ que consiste de todos los automorfismos internos de G . El $\text{Inn}(G)$ es un subgrupo normal del $\text{Aut}(G)$ y el cociente

$$\text{Out}(G) := \frac{\text{Aut}(G)}{\text{Inn}(G)}$$

es el **grupo de automorfismos externos de G** .

Para cualquier $g \in G$, se denotará por $I_g : G \rightarrow G$ al automorfismo interno dado por $h \mapsto ghg^{-1}$.

1.2. Superficies Hiperbólicas

1.2.1. Superficies

Definición 1.12. Una **superficie** es un espacio de Hausdorff, segundo numerable, el cual posee una cubierta abierta $\{U_i\}_{i \in I}$ y una colección $\{\phi_i\}_{i \in I}$ de funciones tales que para cada $i \in I$, $\phi_i : U_i \rightarrow V_i \subset \mathbb{R}^2$ es un homeomorfismo (donde $\{V_i\}_{i \in I}$ es una colección de subconjuntos abiertos de $\mathbb{R}^2$). A la pareja (U_i, ϕ_i) se le llama una **carta coordenada** y la colección de cartas $\{(U_i, \phi_i)\}_{i \in I}$ se le conoce como un **atlas**.

Los siguientes son ejemplos de superficies, la mayoría de estos se usarán frecuentemente a lo largo de esta tesis.

Ejemplo 1.13. No es difícil ver que $\mathbb{R}^2$ con la topología usual es una superficie, el mismo espacio $\mathbb{R}^2$ con la función identidad forman un atlas.

Ejemplo 1.14. La esfera $\mathbb{S}^2 := \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 = 1\}$ con la topología inducida por $\mathbb{R}^3$ es una superficie. En efecto, podemos tomar los entornos $U_1 = \mathbb{S}^2 - \{(0, 0, 1)\}$ y $U_2 = \mathbb{S}^2 - \{(0, 0, -1)\}$ los cuales son homeomorfos a $\mathbb{R}^2$.

Ejemplo 1.15. El disco abierto $\mathring{D}^2 := \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 < 1\}$ es una superficie. De hecho, $\mathring{D}^2$ junto con la inclusión forman un atlas con una sola carta coordenada para $\mathring{D}^2$.

Ejemplo 1.16. Sea $\mathbb{S}^1 := \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 = 1\}$ el círculo unitario. El toro $T = \mathbb{S}^1 \times \mathbb{S}^1$ es una superficie homeomorfa al espacio que resulta de identificar los lados opuestos de un cuadrado (ver Figura 1.1).

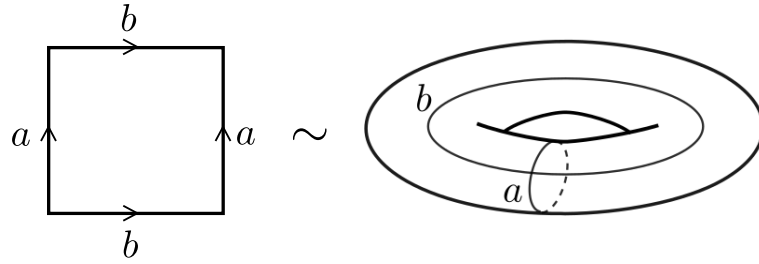


Figura 1.1: El toro como cociente de un cuadrado.

Ejemplo 1.17. Cualquier suma conexa de esferas con toros, o sumas conexas de toros, son superficies. Por ejemplo, en la Figura 1.2 se muestra la suma conexa de dos toros.

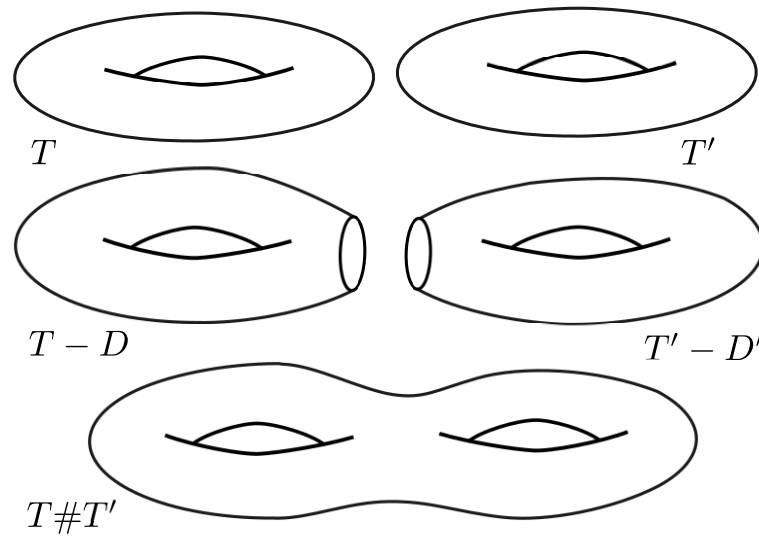


Figura 1.2: Suma conexa de dos toros.

Ejemplo 1.18. Cualquier superficie menos una cantidad finita de puntos es una superficie. Más aún, cualquier superficie menos un conjunto cerrado es una superficie.

La siguiente definición es una generalización de una superficie.

Definición 1.19. Una **superficie con frontera**, es una superficie S en la que además se permite que los conjuntos V_i sean abiertos del semiplano superior $\{(x, y) \in \mathbb{R}^2 \mid y \geq 0\}$. Al conjunto de puntos de S que poseen entornos homeomorfos al semi-

plano superior pero que no poseen entornos homeomorfos a $\mathbb{R}^2$ se llama la **frontera** de S y se denota por ∂S .

Ejemplo 1.20. El disco cerrado $\overline{D^2} := \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 \leq 1\}$ es una superficie cuya frontera es $\mathbb{S}^1$. Cualquier espacio homomorfo a $\overline{D^2}$ se le llamará un disco cerrado, y para relajar la notación, sólo se denotará como D^2 .

Ejemplo 1.21. Cualquier superficie menos una cantidad finita de discos abiertos con cerraduras disjuntas, es una superficie con frontera.

Ejemplo 1.22. Denotaremos por I al intervalo $[0, 1]$. Un anillo A se define como $\mathbb{S}^1 \times I$. Es fácil ver que A es una superficie con frontera: es homomorfo a $\mathbb{S}^2$ menos dos discos abiertos con cerraduras disjuntas.

En lo consiguiente se usará únicamente el termino *superficie* para referirse que puede tener o no frontera.

El siguiente teorema fue demostrado por Möbius a mediados del siglo XIX, el cual es un resultado fundamental para superficies que admiten una triángulación. Más tarde, Radò demostró que cualquier superficie compacta admite una triangulación. Las demostraciones para ambos teoremas se pueden encontrar, por ejemplo, en [14].

Teorema 1.23 (Teorema de clasificación de superficies). *Cualquier superficie orientable y compacta es homeomorfa (de forma única) a la suma conexa de una esfera (de dimensión 2) con $g \geq 0$ toros, removiendo $b \geq 0$ discos abiertos con cerraduras disjuntas. Es decir, el conjunto de superficies compactas y orientables módulo homeomorfismo está en correspondencia biyectiva con el conjunto $\{(g, b) \mid g, b \in \mathbb{Z}, g, b \geq 0\}$.*

El número g en el Teorema 1.23 se le conoce como el **género** de la superficie, el cual indica el número de toros que hay en la suma conexa; el número b es el número de **componentes de frontera** de la superficie. Una manera de obtener una superficie

no compacta a partir de una superficie compacta S es removiendo n puntos en el interior de S ; en este caso, se dice que la superficie resultante tiene n **ponchaduras**.

A menos que se especifique lo contrario, una *superficie* en este trabajo será una superficie orientada, conexa, compacta y posiblemente con ponchaduras (evidentemente, después de ponchar a una superficie compacta, deja de ser compacta). Con esta convención cualquier superficie se puede etiquetar con la tripleta (g, b, n) . Se denotará por $S_{g,n}$ a una superficie de género g con n ponchaduras y frontera vacía; tal superficie es homeomorfa al interior de una superficie compacta con n componentes de frontera. Además, para una superficie de género g , compacta y sin frontera, se abreviará a $S_{g,0}$ como S_g .

La *característica de Euler* de una superficie S se define como

$$\chi(S) := 2 - 2g - (b + n).$$

Es un invariante de la clase de homeomorfismos de S , de esto se sigue que una superficie S está determinada, salvo homeomorfismo, por cualquier terna formada por los números g , b , n , y $\chi(S)$.

A menudo es conveniente pensar a las ponchaduras de una superficies como **puntos marcados**. Esto es, en lugar de eliminar puntos, podemos etiquetarlos. Se supondrá, al menos que se especifique lo contrario, que ninguna superficie tiene puntos marcados.

Si una superficie S es tal que $\chi(S) \leq 0$ y $\partial S = \emptyset$, entonces su cubriente universal $\tilde{S}$ es homeomorfo a $\mathbb{R}^2$ (véase [5]). Posteriormente se verá que, cuando $\chi(S) < 0$, se puede tener mayor provecho de una estructura hiperbólica en $\tilde{S}$.

1.2.2. Superficies Hiperbólicas

La *geometría hiperbólica* es el estudio de las propiedades de los espacios geométricos con *curvatura* constante negativa (ver [13]).

En dimensión 2, las superficies de curvatura constante se clasifican por su curvatura K , teniendo las tres opciones siguientes: positiva, cero o negativa. Dada una superficie S en $\mathbb{R}^3$, si $K > 0$, todos los puntos de la superficie S se encuentran en un lado del plano tangente en cualquier punto; si $K = 0$, siempre existe una línea recta en S contenida en el plano tangente; si $K < 0$, siempre hay puntos de S en ambos lados de cualquier plano tangente. Si el espacio es simplemente conexo y $K > 0$, es una esfera (con $K = 1/\text{radio}$); si $K = 0$, entonces es el plano Euclidiano; si $K < 0$, es el plano hiperbólico, también llamado el espacio hiperbólico de dimensión 2.

Cualquier superficie puede ser dotada de una *estructura geométrica*. Esto significa que se puede encontrar una métrica en la superficie de tal forma que en regiones muy pequeñas se parezcan a alguno de los tres espacios geométricos mencionados anteriormente.

Definición 1.24. Se dice que una superficie S **admite una métrica hiperbólica** si existe una región completa de área finita en S con métrica Riemanniana de curvatura constante -1 donde la frontera de S (si es no vacía) es totalmente geodésico (esto significa que las geodésicas en ∂S son geodésicas en S).

Si S es una superficie sin frontera y tiene una métrica hiperbólica, entonces su cubriente universal $\tilde{S}$ es una superficie Riemanniana simplemente conexa de curvatura constante -1 . De esto se sigue que $\tilde{S}$ es isométrico a $\mathbb{H}$, y S es isométrico a $\mathbb{H}$ módulo una acción por isometrías, libre y propiamente discontinua del $\pi_1(S, p)$. Si S tiene una métrica hiperbólica y su frontera es no vacía, entonces $\tilde{S}$ es isométrico a un subespacio totalmente geodésico de $\mathbb{H}$.

Teorema 1.25. *Sea S una superficie (quizá con ponchaduras o con frontera). Si $\chi(S) < 0$, entonces S admite una métrica hiperbólica.*

Una demostración del teorema anterior se puede encontrar en [5].

Definición 1.26. Una **superficie hiperbólica** es una superficie dotada de una métrica hiperbólica.

Una manera de obtener una métrica hiperbólica en una superficie cerrada S_g de género $g \geq 2$, es construir, para algún $p \in S$, una acción por isometrías, libre y propiamente discontinua de $\pi_1(S, p)$ en $\mathbb{H}$. De la teoría de espacios cubrientes y la clasificación de superficies, el cociente será homeomorfo a S_g . Otra forma de obtener una métrica hiperbólica en S_g , es tomar un $4g$ -gono geodésico en $\mathbb{H}$ (ver Figura 1.3) tal que la suma de sus ángulos interiores sea 2π e identificar sus lados opuestos (tal $4g$ -gono siempre existe; ver, por ejemplo, la Sección 10.4 de [4]). El resultado es una superficie de género g con una métrica hiperbólica y, de acuerdo a la Definición 1.26, su cubriente universal es $\mathbb{H}$.

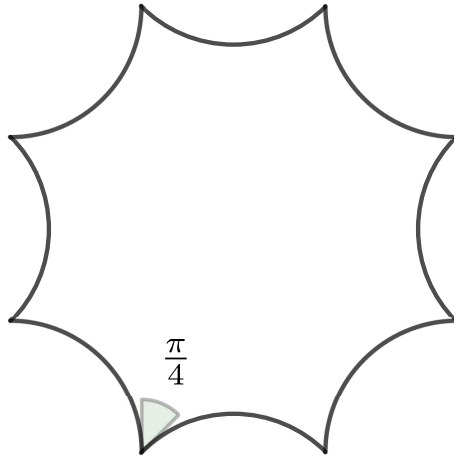


Figura 1.3: 8-gono cuya suma de ángulos interiores es 2π .

Definición 1.27. Sea S una superficie hiperbólica. Una **vecindad de una ponchadura** es un subconjunto cerrado de S homeomorfo a un disco cerrado sin un punto en su interior.

Teorema 1.28. *Si S es una superficie que admite una métrica hiperbólica, entonces el centralizador de cualquier elemento no trivial de $\pi_1(S, p)$ es cíclico. En particular, el centro de $\pi_1(S, p)$ es trivial.*

1.3. Curvas Cerradas Simples

1.3.1. Curvas Cerradas

Definición 1.29. Un **camino cerrado** en una superficie S es una función continua $\alpha : \mathbb{S}^1 \rightarrow S$.

Obsérvese que el camino cerrado es la función α y no la imagen $\alpha([0, 1])$; a la imagen $\alpha([0, 1])$ se le llama una **curva cerrada** de S . En lo consiguiente se identificará a un camino cerrado por su curva cerrada correspondiente.

Definición 1.30. Una curva cerrada es **homotópica a una ponchadura** si es homotópica a la frontera de una vecindad de una ponchadura.

Definición 1.31. Una curva cerrada es **esencial** si no es homotópica a un punto, o a una ponchadura, o a una componente de frontera. Por ejemplo, en la Figura 1.4 las curvas α y β son curvas esenciales, mientras que la curva γ no es esencial.

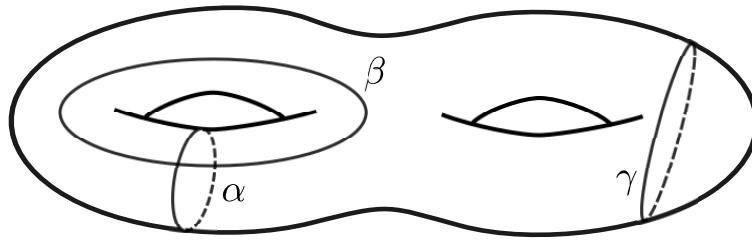


Figura 1.4: α y β son curvas esenciales mientras que γ no lo es.

Curvas cerradas y el grupo fundamental

En el Apéndice B del grupo fundamental se da una definición de camino, el cual es equivalente a la definición de camino cerrado cuando el *origen* y el *final* del camino coinciden. Así, se puede hablar indistintamente de clases de homotopía de curvas cerradas como clases de homotopía de caminos basados en un punto.

Dada una curva cerrada α en una superficie compacta S , esta se puede identificar con un elemento del grupo fundamental $\pi_1(S, p)$ (ver Apéndice B) tomando un camino γ con origen en p y final en algún punto de α . Entonces, la clase de homotopía de $\gamma\alpha\gamma^{-1}$ es un elemento del $\pi_1(S, p)$.

El elemento resultante del $\pi_1(S, p)$ está bien definido salvo conjugación, por ejemplo, en la Figura 1.5, $[\gamma_1 * \alpha * \gamma_1^{-1}]$ es igual a $[\gamma_1 * \tilde{\alpha} * \gamma_2^{-1}][\gamma_2 * \alpha * \gamma_2^{-1}][\gamma_2 * \tilde{\alpha}^{-1} * \gamma_1^{-1}]$, donde $\tilde{\alpha}$ es una curva homotópica a un segmento de α que une al punto final de γ_1 con el punto final de γ_2 . Abusando de la notación, se denotará a este elemento del $\pi_1(S, p)$ por α .

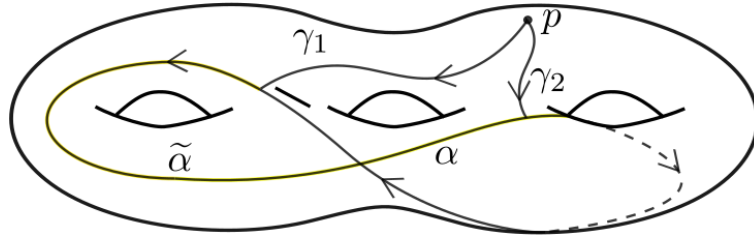


Figura 1.5:

De esta manera es posible identificar la clase de homotopía libre de cualquier curva cerrada y orientada en S con una clase de conjugación del grupo fundamental. Esto deja una correspondencia biyectiva:

$$\left\{ \begin{array}{c} \text{Clases no triviales} \\ \text{de conjugación} \\ \text{en } \pi_1(S, x) \end{array} \right\} \longleftrightarrow \left\{ \begin{array}{c} \text{Clases no triviales de} \\ \text{homotopía libre de curvas} \\ \text{cerradas y orientadas en } S \end{array} \right\}$$

Donde una clase de conjugación (o de homotopía libre) es no trivial si sus elementos son esenciales.

1.3.2. Curvas Cerradas Simples

Definición 1.32. Una curva cerrada $\alpha : \mathbb{S}^1 \rightarrow S$ es **simple** si α es un homeomorfismo a su imagen, en otras palabras, si $\mathbb{S}^1$ está encajado en S mediante α .

Una razón de especial importancia del estudio de las curvas cerradas simples es que éstas se pueden clasificar hasta homeomorfismos de S (ver Sección 1.4). Más aún, es posible estudiar a los homeomorfismos de una superficie vía su acción en curvas cerradas simples.

Antes de ver algunas propiedades de las curvas cerradas simples es conveniente dar las siguientes definiciones.

Definición 1.33. Se dice que un elemento g de un grupo G es **primitivo** si no existe algún $h \in G$, tal que $g = h^k$, con $|k| > 1$.

Nótese que la propiedad de ser primitivo es invariante por automorfismos, y por tanto es un invariante en las clases de conjugación.

Definición 1.34. Una curva cerrada en S es un **múltiplo** si es una función continua $\mathbb{S}^1 \rightarrow S$ que se factoriza mediante la función $\mathbb{S}^1 \xrightarrow{\times n} \mathbb{S}^1$ para $n > 1$. En otras palabras, una curva es un múltiplo si “corre al rededor” de otra curva múltiples veces.

Si una curva cerrada en S es un múltiplo, entonces ningún elemento de su clase de conjugación correspondiente en el $\pi_1(S, p)$ es primitivo.

En la Subsección B.2.2 se da una definición de *levantamiento*, en esta sección es conveniente pensar a un levantamiento de una forma ligeramente distinta.

Definición 1.35. Sea $\rho : \tilde{S} \rightarrow S$ una función cubriente (ver Apéndice B) y α una curva cerrada en S . Un **levantamiento** de α a $\tilde{S}$ se pensará como la imagen de un

levantamiento (en el sentido de la Definición B.33) $\mathbb{R} \rightarrow \tilde{S}$ de la función $\alpha \circ \rho'$, donde $\rho' : \mathbb{R} \rightarrow \mathbb{S}^1$ es la aplicación cubriente usual de $\mathbb{S}^1$; es decir, la imagen de una función $\mathbb{R} \rightarrow \tilde{S}$ tal que el diagrama

$$\begin{array}{ccc} & & \tilde{S} \\ & \nearrow & \downarrow \rho \\ \mathbb{R} & \xrightarrow{\alpha \circ \rho'} & S \end{array}$$

es conmutativo.

Nótese que el sentido de un levantamiento de una curva, en este capítulo, es diferente al sentido usual de un levantamiento de caminos (Definición B.33), el cual es típicamente un subconjunto propio de un levantamiento. Así, por ejemplo, si S es una superficie de género $g \geq 2$, entonces un levantamiento de una curva esencial cerrada simple de S al cubriente universal es una copia de $\mathbb{R}$.

Ahora bien, supóngase que $\tilde{S}$ es el cubriente universal de S y que α es una curva cerrada simple en S que no es un múltiplo no trivial de otra curva cerrada. En este caso, los levantamientos de α a $\tilde{S}$ están en biyección natural con las clases laterales en $\pi_1(S, p)$ del subgrupo cíclico infinito $\langle \alpha \rangle$ (ver Figura 1.6). (Cualquier múltiplo no trivial de α tiene el mismo conjunto de levantamientos que α pero más clases laterales). El grupo $\pi_1(S, p)$ actúa en el conjunto de levantamientos de α por transformaciones de cubierta, y esta acción coincide con la acción izquierda usual de $\pi_1(S, p)$ en las clases laterales de $\langle \alpha \rangle$. El estabilizador del levantamiento correspondiente a la clase lateral $\gamma \langle \alpha \rangle$ es el grupo cíclico $\langle \gamma \alpha \gamma^{-1} \rangle$.

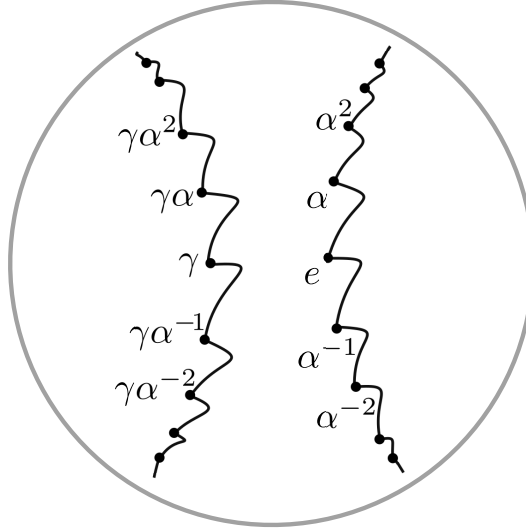


Figura 1.6: La curva de la derecha es el levantamiento de α correspondiente a $\langle\alpha\rangle$ y la curva de la izquierda es levantamiento de α correspondiente a $\gamma\langle\alpha\rangle$.

Cuando S admite una métrica hiperbólica y α es un elemento primitivo del $\pi_1(S, p)$, se tiene la siguiente correspondencia biyectiva

$$\left\{ \begin{array}{c} \text{Elementos de la clase de} \\ \text{conjugación de } \alpha \text{ en } \pi_1(S, p) \end{array} \right\} \longleftrightarrow \left\{ \begin{array}{c} \text{Levantamientos de la} \\ \text{curva cerrada } \alpha \text{ a } S \end{array} \right\}$$

Más precisamente, el levantamiento de la curva α dado por la clase lateral $\gamma\langle\alpha\rangle$ corresponde al elemento $\gamma\alpha\gamma^{-1}$ de la clase de conjugación $[\alpha]$ en $\pi_1(S, p)$. Que esto sea una correspondencia biyectiva es consecuencia del hecho de que, para una superficie hiperbólica S y para cualquier $p \in S$, el centralizador de cualquier elemento de $\pi_1(S, p)$ es cíclico.

Si α es cualquier múltiplo, entonces aún se tiene una correspondencia biyectiva entre los elementos de la clase de conjugación de α y los levantamientos de α . Sin embargo, si α no es primitivo y tampoco múltiplo, entonces hay más levantamientos de α que conjugados. En efecto, si $\alpha = \beta^k$, donde $k > 1$, entonces $\beta\langle\alpha\rangle \neq \langle\alpha\rangle$ mientras

que $\beta\alpha\beta^{-1} = \alpha$.

Representantes geodésicos

Una superficie hiperbólica se caracteriza por tener al plano hiperbólico como cubriente universal Riemanniano (ver Apéndice B). Esto permite probar, por ejemplo, que cada clase de homotopía libre de cualquier curva cerrada contiene una única geodésica.

Definición 1.36. Sea S una superficie hiperbólica. Se dice que una curva α en S es **geodésica** si existe un levantamiento de α a $\mathbb{H}$ que sea una geodésica de $\mathbb{H}$ (ver Definición A.2).

Proposición 1.37. Sea S una superficie hiperbólica. Si α es una curva cerrada en S que no es homotópica a un punto ni a una vecindad de una ponchadura, entonces α es homotópica a una única curva geodésica cerrada γ .

La idea de la demostración es considerar un levantamiento $\tilde{\alpha}$ a $\mathbb{H}$ de α . Como se ha afirmado antes, $\tilde{\alpha}$ está estabilizado por un algún elemento de la clase de conjugación de $\pi_1(S, p)$ correspondiente a α . Se considera la isometría de $\mathbb{H}$ correspondiente a este elemento. Tal isometría resulta ser una isometría hiperbólica (ver Subsección A.1.3), y por tanto tiene un eje de translación, el cual es una geodésica. La proyección de esta geodésica sobre S (vía la función cubriente) resulta ser la curva geodésica cerrada deseada. La demostración completa se puede consultar como la demostración de la Proposición 1.3 de [4].

De la Proposición 1.37 se sigue que para una superficie hiperbólica compacta, se cumple la correspondencia biyectiva siguiente.

$$\left\{ \begin{array}{c} \text{Clases de conjugación} \\ \text{en } \pi_1(S, p) \end{array} \right\} \longleftrightarrow \left\{ \begin{array}{c} \text{Curvas geodésicas cerradas} \\ \text{y orientadas en } S \end{array} \right\}$$

Otra de importancia del estudio de las curvas cerradas simples es porque son una representación de elementos primitivos del $\pi_1(S, p)$.

Proposición 1.38. *Sea α una curva cerrada simple en una superficie S . Si α no se anula homotópicamente, entonces cada elemento de su correspondiente clase de conjugación en $\pi_1(S, p)$ es primitivo.*

La demostración de este resultado se puede consultar como la demostración de la Proposición 1.4 de [4].

Nótese que la Proposición 1.37 es válida para cualquier curva cerrada (no necesariamente simple), en el caso de curvas cerradas simples se tiene además el siguiente resultado.

Proposición 1.39. *Sea S una superficie hiperbólica y α una curva cerrada en S no homotópica a una vecindad de una ponchadura. Llámese γ a la única geodésica en la clase de homotopía libre de α garantizada por la Proposición 1.37. Si α es simple, entonces γ es simple.*

Antes de dar una demostración para la proposición anterior, es preciso definir cuando un levantamiento es simple.

Definición 1.40. Un levantamiento $\tilde{\alpha}$ de una curva α es **simple** si un homeomorfismo a su imagen, es decir, si $\mathbb{R}$ está encajado en $\mathbb{H}$ mediante α .

Demostración. Se partirá del siguiente hecho:

Una curva cerrada β en una superficie hiperbólica S es cerrada si y sólo si se cumplen las siguientes propiedades:

1. *Cada levantamiento de β a $\mathbb{H}$ es simple.*
2. *Ningún par de levantamientos de β se corta.*
3. *β no es un múltiplo (no trivial) de otra curva cerrada.*

Como α es simple, entonces cualesquiera dos levantamientos de α a $\mathbb{H}$ no se cortan. De esto se sigue que para cualesquiera dos levantamientos, sus puntos finales no están ligados en la frontera de $\mathbb{H}$ (ver Definición 4.4). Se tiene que cada levantamiento de γ comparte ambos puntos finales con algún levantamiento de β . Por tanto, los puntos finales de cualesquiera dos levantamientos de γ no están ligados en la frontera de $\mathbb{H}$. Puestos que estos levantamientos son geodésicas, se sigue que estos no se cortan en $\mathbb{H}$. Más aún, por la Proposición 1.38, cualquier elemento de $\pi_1(S, p)$ en la clase de conjugación correspondiente α es primitivo. Lo mismo es cierto para γ , y por esto γ no puede ser un múltiplo. Puesto que las geodésicas en $\mathbb{H}$ son siempre simples, se concluye que γ es simple.

□

1.3.3. Números de Intersección

Definición 1.41. Sea Γ una colección de curvas cerradas simples en una superficie S . Las curvas de Γ están en **posición general** si ninguna curva de Γ es (libremente) homotópica a otra curva de Γ o a sus inversos.

Definición 1.42. Sean α y β un par de curvas cerradas simples, orientadas, y en posición general en S . El **número de intersección algebraico** $\hat{i}(\alpha, \beta)$ se define como la suma de los índices de los puntos de intersección de α y β , donde un punto de intersección es de índice $+1$ cuando la orientación de la intersección coincide con la orientación de S , y -1 de otro modo.

En particular, tiene sentido escribir $\hat{i}(a, b)$ para a y b , las clases de homotopía de curvas cerradas α y β .

Definición 1.43. El **número de intersección geométrico** entre dos clases de homotopía libre a y b de curvas cerradas simples en una superficie S , se define como

el número mínimo de puntos de intersección entre una curva representante de la clase a y una curva representante de la clase b :

$$i(a, b) = \min\{|\alpha \cap \beta| : \alpha \in a \text{ y } \beta \in b\}.$$

En algunas ocasiones, abusando de la notación, se escribirá simplemente $i(\alpha, \beta)$ para denotar al número de intersección entre dos clases de homotopía de curvas cerradas simples α y β .

Nótese que el número de intersección geométrico es simétrico, mientras que el número de intersección algebraico es antisimétrico, es decir, $i(a, b) = i(b, a)$ mientras que $\hat{i}(a, b) = -\hat{i}(b, a)$.

Obsérvese que $i(a, a) = 0$ para cualquier clase de homotopía libre de una curva cerrada simple a . Si α separa a S en dos componentes, entonces para cualquier β se tiene que $\hat{i}(\alpha, \beta) = 0$ y $i(\alpha, \beta)$ es par. En general, i y $\hat{i}$ tienen la misma paridad, e $\hat{i}$ es menor o igual i .

1.3.4. Criterio del Bígono

Nótese que en la definición de número de intersección geométrico se utilizó el mínimo, pues siempre se pueden encontrar dos representantes α y β de las clases de homotopía de a y b , tales que $i(a, b) = |\alpha \cap \beta|$. Esto motiva la siguiente definición.

Definición 1.44. Sean α y β dos curvas cerradas en una superficie S . Supóngase que a y b son la clases de homotopía de α y β respectivamente. Se dice que α y β están en **posición mínima** si $i(a, b) = |\alpha \cap \beta|$.

Definición 1.45. Sean α y β dos curvas cerradas simples en posición general en una superficie S . Se dice que α y β forman un **bígono** si hay un disco encajado en S cuya

frontera es la unión de un arco de α y un arco de β intersecados en exactamente dos puntos.

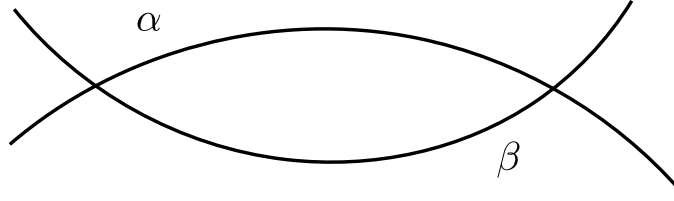


Figura 1.7: Las curvas α y β forma un bígono.

Lema 1.46. *Si dos curvas cerradas simples α y β en posición general en una superficie S no forman un bígono, entonces en el cubriente universal de S , cualquier par de levantamientos $\tilde{\alpha}$ y $\tilde{\beta}$ de α y β tienen a lo más un punto de intersección.*

Proposición 1.47 (Criterio del Bígono). *Dos curvas cerradas simples en posición general en una superficie S están en posición mínima si y sólo si no forman un bígono.*

Corolario 1.48. *Cualesquiera dos curvas cerradas simples en posición general que se intersectan una sola vez están en posición mínima.*

Geodésicas están en posición mínima Obsérvese que si dos segmentos geodésicos sobre una superficie S formasen un bígono, tal bígono por ser simplemente conexo pudiese levantarse al cubriente universal $\mathbb{H}$ de S , en donde se obtendría una región bordeada por dos geodésicas distintas. Esto contradice el hecho de que las geodésicas entre cualesquiera dos puntos en $\mathbb{H}$ son únicas. Por tanto, se tiene lo siguiente:

Corolario 1.49. *Dos geodésicas cerradas simples distintas en una superficie hiperbólica están en posición mínima.*

Si el lector esta interesado en conocer las demostraciones del Lema 1.46, la Proposición 1.47, el Corolario 1.48 y el Corolario 1.49, éstas se pueden consultar, por ejemplo, en la Subsección 1.2.4 de [4].

La Proposición 1.37 junto con el Corolario 1.49 implican que, dados cualquier número de clases distintas de homotopía de curvas cerradas simples esenciales en S , se puede elegir un representante particular de cada clase (por ejemplo, una geodésica) tal que cada par de curvas están en posición mínima.

1.4. Principio De Cambio de Coordenadas

En esta sección se describe una técnica frecuentemente utilizada en la teoría del grupo de modular de Teichmüller, conocida como el **principio de cambio de coordenadas**, el cual es análogo al cambio de base en espacio vectorial.

1.4.1. Clasificación de curvas cerradas simples.

Dada una curva cerrada simple α en una superficie S , la superficie obtenida por **cortar** a S a lo largo de α es una superficie compacta S_α equipada con un homeomorfismo h entre dos de sus componentes de frontera tal que

- i) el cociente $S_\alpha/(x \sim h(x))$ es homeomorfo a S , y
- ii) la imagen de estas distinguidas componentes de frontera bajo la función cociente es α .

Tiene sentido también cortar una superficie con frontera o puntos marcados a lo largo de un arco simple propio; la definición es análoga. Similarmente, se puede cortar a través de una colección finita de curvas y arcos.

Definición 1.50. Una curva cerrada simple α (o un arco simple) en una superficie S es **separadora**, si al cortar sobre α , la superficie resultante S_α es disconexa. De lo contrario, si S_α es conexa, se dice que α es **no separadora**.

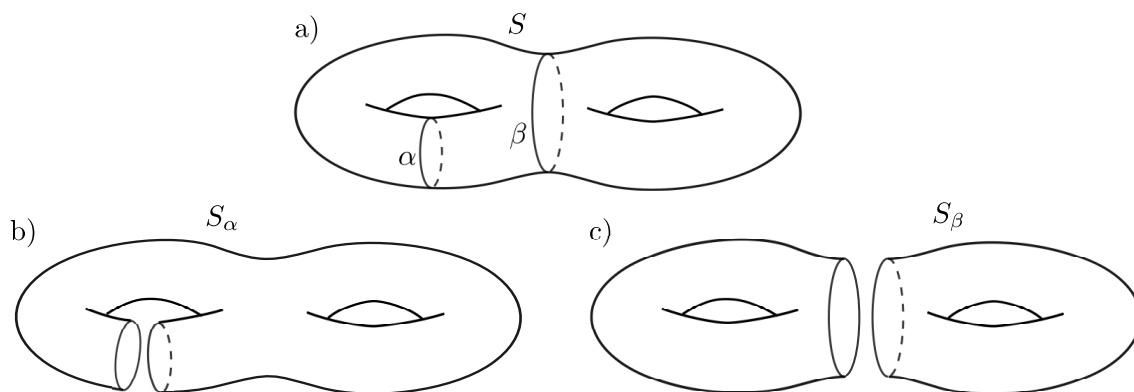


Figura 1.8: En la superficie S , α no es una curva separadora pero β sí lo es.

Proposición 1.51. Si α y β son cualesquiera dos curvas cerradas simples no separadoras en una superficie S , entonces existe un homeomorfismo $\phi : S \rightarrow S$ tal que $\phi(\alpha) = \beta$.

En otras palabras, salvo homeomorfismos, no hay más que una curva cerrada simple no separadora en una superficie S . La demostración de esta proposición se sigue del principio de cambio de coordenadas: Si α y β son dos curvas no separadoras en una superficie S , entonces las superficies S_α y S_β tienen cada una dos componentes de frontera correspondientes a α y β . Puesto que S_α y S_β tienen la misma característica de Euler, el mismo número de componentes de frontera y mismo número de ponchaduras, se sigue que S_α es homeomorfo a S_β . Podemos elegir un homeomorfismo $S_\alpha \rightarrow S_\beta$ que respete las relaciones de equivalencia en las componentes de frontera. Tal homeomorfismo induce un homeomorfismo $\phi : S \rightarrow S$ que manda a α a β . El homeomorfismo ϕ se piensa como un cambio de coordenadas, puesto que transforma la curva α a la curva β .

Supóngase que α_1 y β_2 son dos curvas cerradas simples en una superficie S que se intersectan una vez. Sea S_{α_1} la superficie obtenida al cortar a S a lo largo de α_1 . S_{α_1} tiene dos componentes de frontera correspondientes a los lados de α_1 . La imagen de β_1 es un arco simple que conecta estas componentes de frontera. Podemos cortar a S_{α_1} a través de este arco para obtener una superficie $(S_{\alpha_1})_{\beta_1}$. Esta última superficie tiene solamente una componente de frontera que es subdivida naturalmente en cuatro arcos (dos provenientes de α_1 y dos provenientes de β_1). La relación de equivalencia en la definición de cortar una superficie, identifica estos arcos con el fin de recuperar a la superficie S con sus curvas α_1 y β_2 .

Si α_2 y β_2 son otra colección de curvas cerradas simples que se intersectan una vez, entonces, análogamente hay una superficie cortada $(S_{\alpha_2})_{\beta_2}$. Por la clasificación de superficies, $(S_{\alpha_2})_{\beta_2}$ es homeomorfo a $(S_{\alpha_1})_{\beta_1}$, más aún existe un homeomorfismo que preserva las clases de equivalencia en la frontera. Cualquiera de estos homeomorfismos descende a un homeomorfismo de S que manda al par $\{\alpha_1, \beta_1\}$ al par $\{\alpha_2, \beta_2\}$.

Un poco más general de caso anterior, es cuando se tiene una *cadena* de curvas cerradas simples.

Definición 1.52 (Cadena). Una **cadena** de curvas cerradas simples en una superficie S es una secuencia $\alpha_1, \dots, \alpha_k$ (de curvas cerradas simples) en S con la propiedad de que $i(\alpha_i, \alpha_{i+1}) = 1$ para cada i y $i(\alpha_i, \alpha_j) = 0$ siempre que $|i - j| > 1$. Una cadena es *noseparadora* si la unión de las curvas no separa a la superficie.

Teorema 1.53. Sean dos cadenas $\alpha_1, \dots, \alpha_m$ y $\beta_1, \dots, \beta_n$ no separadoras de curvas cerradas simples en una superficie S , si $m = n$ entonces existe un homeomorfismo $\phi : S \rightarrow S$ tal que $\phi(\alpha_i) = \beta_i$.

En la Sección 1.3.3 de [4] se puede consultar un bosquejo de la demostración del Teorema 1.53, en donde se encuentra como el ejemplo 6. Además, se puede probar que cada cadena en S_g de longitud par es no separadora, y por tanto tales cadenas

deben ser topológicamente equivalentes.

Para saber más sobre el principio de cambio de coordenadas y las demostraciones de las proposiciones (teoremas, corolarios, etc.) de esta subsección, el lector puede consultar la Subsección 1.3 de [4].

Todas los teoremas, proposiciones y corolarios presentados en este capítulo, aunque en su mayoría no se mencionan, se utilizan implícitamente en la demostración del Teorema de Dehn-Nielsen-Baer (Capítulo 5) para caracterizar propiedades topológicas de curvas cerradas simples.

Capítulo 2

Grupo Modular de Teichmüller

Este capítulo es una introducción al grupo modular de Teichmüller. Se inicia definiendo la relación de isotopía en el conjunto de curvas cerradas de una superficie S , la cual se utiliza para definir las clases de isotopías de homeomorfismos de S . Posteriormente se da la definición del grupo modular de Teichmüller y se incluyen algunos ejemplos básicos. Finalmente se presenta el método de Alexander, el cual provee una respuesta a una pregunta que surge naturalmente al estudiar al grupo modular de Teichmüller: ¿Cómo saber si dos homeomorfismos son isotópicos o no?, o equivalentemente, ¿cómo probar si un homeomorfismo es o no isotópico a la identidad?

2.1. Clases de Isotopía

Isotopía de Curvas Cerradas Simples

Definición 2.1. Sean $\alpha, \beta : S^1 \rightarrow S$ dos curvas cerradas simples en una superficie S . Se dice que α y β son **isotópicos** si existe una homotopía (ver Apéndice B.1.2)

$$F : S^1 \times I \rightarrow S$$

entre α y β con la propiedad de que la curva $F(S^1 \times \{t\})$ es simple para cada $t \in I$. En este caso, la homotopía F se llama **isotopía** entre α y β .

Si una homotopía entre dos curvas existe, en general no es una isotopía. El siguiente resultado, originalmente demostrado por Baer, dice cuándo una homotopía es una isotopía.

Lema 2.2. *Sean α y β dos curvas cerradas simples esenciales en una superficie S . Entonces α y β son isotópicos si y sólo si α y β son homotópicos.*

Una implicación queda evidente puesto que por definición una isotopía es una homotopía, mientras que la otra se puede encontrar como la demostración de la Proposición 1.10 de [4].

Isotopía de homeomorfismos

Definición 2.3. Sea S una superficie y sean $f_0, f_1 : S \rightarrow S$ dos homeomorfismos. Se dice que f_0 y f_1 son **isotópicos** si existe una homotopía $F : I \times S \rightarrow S$ entre f_0 y f_1 tal que la función $f_t : S \rightarrow S$ dada por $f_t(x) = F(t, x)$ es un homeomorfismo para cada $t \in I$. La función F se llama **isotopía** de S entre f_0 y f_1 .

Evidentemente la relación de isotopía, ya sea en el conjunto de curvas cerradas simples de S o en el grupo de homeomorfismos de S , es una relación de equivalencia. A las clases de equivalencia bajo esta relación se les llaman **clases de isotopía**, y dependiendo del contexto, se tratará de curvas cerradas simples u homeomorfismos.

Análogo a la Proposición 2.1, bajo las hipótesis adecuadas, se puede garantizar que si dos homeomorfismos son homotópicos entonces son isotópicos:

Proposición 2.4. *Sea S una superficie compacta y sean f y g dos homeomorfismos homotópicos de S . Si S no es homeomorfo a un disco cerrado D^2 o un anillo cerrado A , entonces f y g son isotópicos. Más aún, si f y g son homotópicos relativos a ∂S , entonces f y g son isotópicos relativos a ∂S .*

En particular, en cualquier superficie cerrada de género $g \geq 0$, homotopía implica isotopía.

En el caso del disco cerrado D^2 y el anillo cerrado A se pueden encontrar ejemplos de homeomorfismos homotópicos que no son isotópicos (ver, por ejemplo, [4] página 41). De hecho, para cada caso existe un único ejemplo, y se trata de homeomorfismos que no preservan la orientación, los cuales son homotópicos a la identidad pero no isotópicos. Los detalles de la demostración de la Proposición 2.4 se puede consultar, por ejemplo, en [3].

En virtud del Lema 2.2, en superficies compactas, hablar de homotopía es equivalente a hablar de isotopía.

Extensión de Isotopías

Dada una isotopía entre dos curvas cerradas simples en una superficie S , es útil extenderla a una isotopía de S , el cual es llamado una **isotopía ambiente** de S .

Proposición 2.5. *Sea S una superficie. Si $F : S^1 \times I \rightarrow S$ es una isotopía suave de curvas cerradas simples, entonces existe una isotopía $H : S \times I \rightarrow S$ tal que $H(s, 0) = s$ para cada $s \in S$, y $H(F(s, 0), t) = F(s, t)$ para todo $s \in S$ y todo $t \in I$.*

La Proposición 2.5 es un resultado estándar de topología diferencial. Supóngase que las dos curvas son disjuntas. Para construir la isotopía, se inicia definiendo un campo vectorial suave que sea compatible en una vecindad del anillo determinado por las dos curvas y que lleva una curva a la otra. Se obtiene la isotopía de la superficie extendiendo este campo vectorial a S y luego integrándolo. El lector interesado en los detalles puede encontrar más información, por ejemplo, en [7].

2.2. Grupo Modular de Teichmüller

El grupo modular de Teichmüller fue estudiado por primera vez en 1920 por el matemático alemán Max Dehn [10]. Existen distintas maneras equivalentes de definir el grupo modular de Teichmüller (ver, por ejemplo, [4]). Para los fines de este trabajo es conveniente definirlo de la forma siguiente:

Definición 2.6 (Grupo Modular de Teichmüller). Sea S una superficie de género $g \geq 0$, $b \geq 0$ componentes de frontera y $n \geq 0$ ponchaduras. Se denota por $\text{Homeo}^+(S, \partial S)$ al grupo de homeomorfismos de S que preservan la orientación, fijan el conjunto de ponchaduras, y puntualmente fijan la frontera de S . El **grupo modular de Teichmüller de S** , denotado por $\text{Mod}(S)$, se define como $\text{Homeo}^+(S, \partial S)$ módulo isotopía, es decir, $\text{Mod}(S)$ es el grupo de clases de isotopía de homeomorfismos de S que preservan la orientación, en donde las isotopías fijan al conjunto de ponchaduras y puntualmente fijan la frontera de S .

En esta tesis, para facilitar la lectura, al grupo de modular de Teichmüller de S se le llamará únicamente grupo modular de S .

En el caso de una superficie sin frontera (ni puntos marcados o pochaduras), $\text{Mod}(S)$ es simplemente el grupo de clases de isotopía de homeomorfismos de S que preservan la orientación.

Involución Hiperelíptica

Para una superficie S_g de género $g \geq 0$, un ejemplo importante de un elemento no trivial del grupo $\text{Mod}(S_g)$, es la *involución hiperelíptica*, la cual está dada por una rotación de un ángulo π respecto a un eje de simetría de S_g . Este es un elemento de orden dos del grupo modular $\text{Mod}(S_g)$.

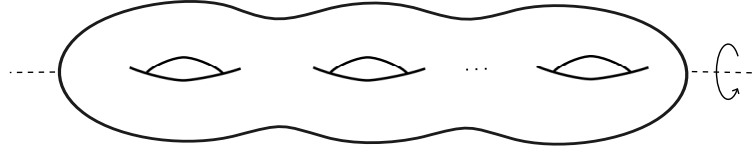


Figura 2.1: La rotación de ángulo π respecto al eje indicado es una involución hiper-elíptica.

Ejemplos del $\text{Mod}(S)$

El primer ejemplo que se verá es el grupo modular de un disco cerrado D^2 .

Lema 2.7 (Lema de Alexander). *El grupo $\text{Mod}(D^2)$ es trivial.*

En otras palabras, el Lema 2.7 declara que dado cualquier homeomorfismo ϕ de D^2 que es la identidad en sobre la frontera ∂D^2 , existe una isotopía de ϕ a la identidad a través de homeomorfismos que son la identidad en ∂D^2 .

Demostración. Identifíquese a D^2 con el disco unitario en $\mathbb{R}^2$. Sea $\phi : D^2 \rightarrow D^2$ un homeomorfismo tal que $\phi|_{\partial D^2}$ es igual a la identidad. Sin pérdida de generalidad, podemos asumir que ϕ deja al origen fijo. Defínase

$$F(x, t) := \begin{cases} (1-t)\phi\left(\frac{x}{1-t}\right) & \text{si } 0 \leq |x| \leq 1-t, \\ x & \text{si } 1-t \leq |x| \leq 1 \end{cases}$$

para $0 \leq t < 1$, y defínase a $F(x, 1)$ como la función identidad de D^2 . La función F es una isotopía de ϕ a la identidad. $\square$

La demostración del Lema 2.7 también es válida cuando se reemplaza a D^2 por un disco con una ponchadura en su interior (identifique la ponchadura/o el punto marcado con el origen), por tanto también se cumple la siguiente proposición:

Proposición 2.8. El grupo modular de un disco con una ponchadura en su interior es trivial.

Existen otros dos grupo modulares $\text{Mod}(S_{g,n})$ que también son triviales, a saber, $\text{Mod}(S_{0,1})$ y $\text{Mod}(S^2)$. Para el primero, se puede identificar a $S_{0,1}$ con $\mathbb{R}^2$ y usar el hecho de que cualquier homeomorfismo $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ que preserva la orientación es homotópico a la identidad vía la homotopía dada por $F(x, t) = (1 - t)x + tf(x)$. Para S^2 , recuérdese que cualquier homeomorfismo de S^2 deja fijo algún punto, así cualquier homeomorfismo puede ser modificado mediante una isotopía que fija un punto, aplicando el ejemplo anterior se llega al resultado.

Ejemplo 2.9. Sea Σ_3 el grupo de permutaciones de tres elementos. Se tiene que

$$\text{Mod}(S_{0,3}) \cong \Sigma_3.$$

Para ver esto es conveniente pensar a $S_{0,3}$ como una esfera con tres puntos marcados (en lugar de tres ponchaduras). El isomorfismo $S_{0,3} \rightarrow \Sigma_3$ está dado por la acción del $\text{Mod}(S_{0,3})$ sobre el conjunto de puntos marcados de $S_{3,0}$.

Ejemplo 2.10. Análogamente se tiene que hay un isomorfismo $\text{Mod}(S_{0,2}) \rightarrow \mathbb{Z}/2\mathbb{Z}$ dado por la acción en los dos puntos marcados.

Ejemplo 2.11. El grupo modular de anillo cerrado A es $\text{Mod}(A) \cong \mathbb{Z}$.

Ejemplo 2.12. El grupo modular del toro T^2 es isomorfo al grupo especial lineal $\text{SL}(2, \mathbb{Z})$. Una demostración de este hecho se encuentra en [4] como la demostración del Teorema 2.5, en tal demostración se utilizan resultados de la teoría de grupos de homología. Si el lector está interesado en grupos de homología, puede consultar, por ejemplo, [6].

2.3. Método de Alexander

Una pregunta que surge naturalmente al estudiar la teoría del grupo modular es: ¿Cuándo dos elementos de este grupo son iguales?, esto es equivalente a preguntar

¿cuándo un homeomorfismo es isotópico a la identidad? Se termina esta sección enunciando el método de Alexander, el cual provee de una respuesta a las preguntas anteriores.

En resumen, el método de Alexander declara que, para cualquier superficie S , un elemento del $\text{Mod}(S)$ a menudo está determinado por su acción en una colección adecuada de curvas y arcos en S . Por tanto, se tiene una manera concreta para determinar cuando dos homeomorfismos $f, g \in \text{Homeo}^+(S, \partial S)$ representan el mismo elemento del $\text{Mod}(S)$.

En general no es cierto que si un homeomorfismo de una superficie S fija una colección de curvas y arcos que separan a S en discos, entonces este es un representante del elemento trivial del $\text{Mod}(S)$.

Antes de enunciar el método de Alexander es conveniente dar la siguiente definición:

Definición 2.13. Sea S una superficie (que posiblemente tiene puntos marcados). Se dice que una colección $\{\gamma_i\}$ de curvas y arcos **llena** a S , si la superficie resultante al cortar a S a través de todas las γ_i es una unión disjunta de discos y discos con un punto marcado en su interior.

Teorema 2.14 (Método de Alexander). *Sea S una superficie compacta, posiblemente con puntos marcados. Supóngase que $\phi : S \rightarrow S$ es un homeomorfismo que preserve la orientación y que fija la frontera de S (en el caso de que S tenga frontera). Sean $\gamma_1, \gamma_2, \dots, \gamma_n$ curvas esenciales cerradas simples y arcos propios en S con las siguientes propiedades.*

1. *Cualesquiera dos curvas γ_i no son homotópicas.*
2. *Cualquier par de curvas γ_i están en posición mínima.*
3. *Para cualesquiera i, j, k distintos, al menos una de las intersecciones $\gamma_i \cap \gamma_j$, $\gamma_i \cap \gamma_k$ o $\gamma_j \cap \gamma_k$ es vacía.*

4. *El conjunto de las curvas $\{\gamma_1, \dots, \gamma_n\}$ llena a S .*

Si ϕ fija las clases de isotopía con orientación de cada γ_i , entonces ϕ es isotópico a la identidad. De otro modo, hay una potencia no trivial de ϕ que es isotópica a la identidad.

La demostración del Teorema 2.3 se puede consultar en [4], en donde aparece como la demostración de la Proposición 2.8.

El poder del método de Alexander radica en que convierte el cálculo de un elemento del grupo de modular en un problema combinatorio finito. Este método se usa frecuentemente, por ejemplo, para calcular el centro del grupo modular (ver Sección 3 de [4]), para demostrar que el problema de la palabra es soluble en $\text{Mod}(S)$ (ver Capítulo 8 de [4]). En este trabajo se usa para demostrar el teorema de Dehn-Nielsen-Baer (Capítulo 5).

Capítulo 3

Lema de Švarc-Milnor

Una de las filosofías del Programa de Erlangen es estudiar a la geometría de un objeto mediante las propiedades que se preservan ante un grupo de transformaciones. Inversamente, uno de los objetivos de la teoría geométrica de grupos es estudiar a los grupos como objetos geométricos. Una manera de hacerlo es induciendo una métrica al grupo, la métrica de las palabras respecto a un conjunto generador. Desafortunadamente, en general, esta métrica depende del conjunto generador seleccionado.

Usando el concepto de cuasi-isometría, la geometría a gran escala permite tener una noción de la geometría de un grupo independiente del conjunto generador seleccionado.

En este capítulo, primero se presentan algunas generalidades de isometrías y cuasi-isometrías. Como segundo paso se aplican estos conceptos a los grupos finitamente generados. Finalmente se demuestra el lema de Švarc-Milnor, el cual es la llave para unir a la teoría geométrica de grupos con la geometría usual.

3.1. Cuasi-isometrías

3.1.1. Encajes isométricos

Definición 3.1. Una función $f : X \rightarrow Y$ entre espacios métricos (X, d_X) y (Y, d_Y) es un **encaje isométrico** si

$$d_Y(f(x), f(x')) = d_X(x, x')$$

para todo $x, x' \in X$. La función f es una **isometría** si es un encaje isométrico y si existe otro encaje isométrico $g : Y \rightarrow X$ tal que

$$f \circ g = \text{id}_Y \text{ y } g \circ f = \text{id}_X.$$

Se dice que dos espacios métricos son **isométricos** si existe una isometría entre ellos.

Los encajes isométricos son inyectivos y continuos; por tanto, cada isometría es un homeomorfismo respecto a las topologías inducidas por las métricas; más aún, el conjunto de todas las isometrías de un espacio métrico X en sí mismo, forman un grupo con la composición usual de funciones.

3.1.2. Encajes Cuasi-Isométricos

Definición 3.2. Sea $f : X \rightarrow Y$ una función entre espacios métricos. Una función $f' : X \rightarrow Y$ **está a distancia finita de f** si existe una $c \in \mathbb{R}_{\geq 0}$ tal que

$$d_Y(f(x), f'(x)) \leq c$$

para todo $x \in X$; donde d_Y denota la métrica en Y .

Definición 3.3. Una función $f : X \rightarrow Y$ entre espacios métricos (X, d_X) y (Y, d_Y) es un *encaje cuasi-isométrico*, si existen constantes $c \in \mathbb{R}_{\geq 1}$ y $b \in \mathbb{R}_{\geq 0}$ tales que

$$\frac{1}{c} \cdot d_X(x, x') - b \leq d_Y(f(x), f(x')) \leq c \cdot d_X(x, x') + b,$$

para todo $x, x' \in X$. La función f es una **cuasi-isometría** si es un encaje cuasi-isométrico, para el cual existe una **cuasi-inversa** $g : Y \rightarrow X$ tal que g es un encaje cuasi-isométrico, es decir, $g : Y \rightarrow X$ es un encaje cuasi-isométrico tal que $g \circ f$ está a distancia finita de id_X y $f \circ g$ está a distancia finita de id_Y . Se dice que dos espacios métricos son **cuasi-isométricos** si existe una cuasi-isometría entre ellos; esto se denota por $X \sim_{QI} Y$.

Si $f : X \rightarrow Y$ es un encaje cuasi-isométrico con constantes c, b (como en la Definición 3.3), se dirá que f es un *encaje (c, b) – cuasi – isométrico*, y si es una cuasi-isometría se dirá que es una *(c, b) – cuasi – isometría*.

A diferencia de los encajes isométricos los encajes cuasi-isométricos no necesariamente son inyectivos o continuos.

Definición 3.4. Sea (X, d_X) un espacio métrico. Un subconjunto $Y \subset X$ es **cuasi-denso** en X si existe una constante $c \in \mathbb{R}_{>0}$ tal que para cada $x \in X$ se puede encontrar una $y \in Y$ tal que

$$d_X(y, x) \leq c.$$

Se dirá que una función $f : X \rightarrow Y$ entre espacios métricos tiene **imagen cuasi-densa** si su imagen $f(X)$ es cuasi-densa en Y .

Proposición 3.5 (Caracterización de una cuasi-isometría). *Una función $f : X \rightarrow Y$ entre espacios métricos (X, d_X) y (Y, d_Y) es una cuasi-isometría si y sólo si es un encaje cuasi-isométrico con imagen cuasi-densa.*

Demostración. Si $f : X \rightarrow Y$ es una cuasi-isometría, por definición f tiene una cuasi-inversa que es un encaje cuasi-isométrico $g : Y \rightarrow X$ tal que $g \circ f$ está a distancia finita de id_X y $f \circ g$ tiene distancia finita de id_Y . El hecho de que $f \circ g$ esté a distancia finita de la identidad id_Y implica que existe una constante $c \in \mathbb{R}_{>0}$ tal que para todo $y \in Y$

$$d_Y(f \circ g(y), y) \leq c;$$

en particular, f tiene imagen cuasi-densa.

Ahora bien, supóngase que $f : X \rightarrow Y$ es un encaje cuasi-isométrico con imagen cuasi-densa. Usando el axioma de la elección, se probará que existe una cuasi-inversa para f .

Dado que f es un encaje cuasi-isométrico con imagen cuasi-densa, existe una constante $c \in \mathbb{R}_{>0}$ (por ejemplo, el máximo de las constantes mencionadas en la Definición 3.3 y en la Definición 3.4) tal que para todo $x, x' \in X$ se satisface que

$$\frac{1}{c} \cdot d_X(x, x') - c \leq d_Y(f(x), f(x')) \leq c \cdot d_X(x, x') + c,$$

y que para todo $y \in Y$ existe $x \in X$ tal que

$$d_Y(f(x), y) < c.$$

Por el axioma de la elección, existe una función

$$\begin{aligned} g : Y &\longrightarrow X \\ y &\longmapsto x_y \end{aligned}$$

tal que $d_Y(f(x_y), y) \leq c$ se satisface para todo $y \in Y$.

Por construcción, para todo $y \in Y$ se cumple que

$$d_Y(f \circ g(y), y) = d_Y(f(x_y), y) \leq c.$$

Inversamente, para todo $x \in X$ se tiene (usando el hecho de que f es un encaje cuasi-isométrico) que

$$d_X(g \circ f(x), x) = d_X(x_{f(x)}, x) \leq c \cdot d_Y(f(x_{f(x)}), f(x)) + c^2 \leq c \cdot c + c^2 = 2 \cdot c^2.$$

De lo anterior se deduce que $f \circ g$ está a distancia finita de la función identidad en Y , y que $g \circ f$ está a distancia finita de la función identidad en X .

Ahora se probará que g también es un encaje cuasi-isométrico. Si $y, y' \in Y$, entonces

$$\begin{aligned} d_X(g(y), g(y')) &= d_X(x_y, x_{y'}) \\ &\leq c \cdot d_Y(f(x_y), f(x_{y'})) + c^2 \\ &\leq c \cdot (d_Y(f(x_y), y) + d_Y(y, y') + d_Y(f(x_{y'}), y')) + c^2 \\ &\leq c \cdot (d_Y(y, y') + 2 \cdot c) + c^2 \\ &= c \cdot d_Y(y, y') + 3 \cdot c^2 \end{aligned}$$

y

$$\begin{aligned} d_X(g(y), g(y')) &= d_X(x_y, x_{y'}) \\ &\geq \frac{1}{c} \cdot d_Y(f(x_y), f(x_{y'})) - 1 \\ &\geq \frac{1}{c} \cdot (d_Y(y, y') - d_Y(f(x_y), y) - d_Y(f(x_{y'}), y')) - 1 \\ &\geq \frac{1}{c} \cdot d_Y(y, y') - \frac{2 \cdot c}{c} - 1. \end{aligned}$$

□

El mismo argumento demuestra que las cuasi-ivasas de los encajes cuasi-isométricos también son encajes cuasi-isométricos.

Proposición 3.6 (Propiedades de Encajes Cuasi-Isométricos).

1. *Cualquier función que está a distancia finita de un encaje cuasi-isométrico es también un encaje cuasi-isométrico.*
2. *Cualquier función que está a distancia finita de una cuasi-isometría es una cuasi-isometría.*
3. *Sean X, Y, Z espacios métricos y sean $f, f' : X \rightarrow Y$ que están a distancia finita entre ellas.*
 - a) *Si $g : Z \rightarrow X$ es una función, entonces $f \circ g$ y $f' \circ g$ están a distancia finita entre ellas.*
 - b) *Si $g : Y \rightarrow Z$ es un encaje cuasi-isométrico, entonces $g \circ f$ y $g \circ f'$ también están a distancia finita entre ellas.*
4. *Composiciones de encajes cuasi-isométricos son encajes cuasi-isométricos.*
5. *Composiciones de cuasi-isometrías son cuasi-isometrías.*

La Definición 3.2 induce una relación de equivalencia en el conjunto de encajes cuasi-isométricos de un espacio métrico X a un espacio métrico Y ; más aún, la Proposición 3.6 garantiza que esta relación es compatible con la composición usual de funciones. De esto, se puede construir un grupo a partir del conjunto de cuasi-isometrías de un espacio en si mismo módulo distancia finita. Este grupo se le conoce como el grupo de cuasi-isometrías de X .

3.1.3. Aplicaciones en Grupos Finitamente Generados

Definición 3.7. Sea G un grupo y $A \subset G$ un conjunto generador de G . La métrica de las palabras d_A sobre G respecto a S se define como

$$d_A(g, h) = \min\{n \in \mathbb{N} \mid \text{Existen } a_1, \dots, a_n \in A \cup A^{-1}, g^{-1}h = a_1 \cdots a_n\}$$

para todo $g, h \in G$.

La métrica de las palabras es en efecto una métrica en G , que además es invariante respecto a la multiplicación por la izquierda, es decir,

$$d_A(lg, lh) = d(g, h),$$

para todo $g, h, l \in G$. Esto último es consecuencia de que $(lg)^{-1}(lh) = g^{-1}h$.

Es importante notar que la métrica de las palabras d_A en un grupo G depende de la elección del conjunto generador A . La noción de cuasi-isometría permite relacionar propiedades geométricas a gran escala que no dependen de la elección del conjunto generador, siempre que este sea finito.

Teorema 3.8. Sean A y B subconjuntos generadores de un grupo G . Si A y B son finitos, entonces se cumple que

1. La identidad $Id_G : G \rightarrow G$ es una cuasi-isometría entre (G, d_A) y (G, d_B) .
2. En particular, cualquier espacio métrico (X, d) cuasi-isométrico a (G, d_A) también es cuasi-isométrico a (G, d_B) .

Demostración. La segunda parte de la proposición anterior se deriva directamente de la primera, pues la composición de cuasi-isometría es una cuasi-isometría (Proposición 3.6).

Así, es suficiente con demostrar la primera parte. Usando la caracterización de cuasi-isometrías (Proposición 3.5) basta probar que $Id_G : G \rightarrow G$ es un encaje cuasi-isométrico con imagen cuasi-densa. Evidentemente Id_G tiene imagen cuasi-densa, así sólo se deben encontrar constantes $c > 0$ y $b \geq 0$, tales que para cualesquiera $g, h \in G$, se cumpla que

$$\frac{1}{c} \cdot d_{s'}(g, h) - b \leq d_s(g, h) \leq c \cdot d_{s'}(g, h) + b.$$

Nótese que para cada $a \in A \cup A^{-1}$, existe $k \in \mathbb{N}$ y ciertos $b_1, \dots, b_k \in B \cup B^{-1}$ tales que $a = b_1 \cdots b_k$. Además, dado que A es finito, el conjunto $\{d_B(e, a) \mid a \in A \cup A^{-1}\}$ tiene un máximo finito. Tómese a

$$c_1 = \max_{a \in A \cup A^{-1}} d_B(e, a).$$

Sean g y h dos elementos arbitrarios de G y sea $n = d_A(g, h)$. Dado que A es un conjunto generador, se tiene que $g^{-1}h$ es igual $a_1 \cdots a_n$ para ciertos $a_1, \dots, a_n \in A \cup A^{-1}$, lo que implica que $h = ga_1 \cdots a_n$. Utilizando la desigualdad del triángulo y que d_B es invariante por la multiplicación por la izquierda, se obtiene que

$$\begin{aligned} d_B(g, h) &= d_B(g, ga_1 \cdots a_n) \\ &\leq d_B(g, ga_1) + d_B(ga_1, ga_1a_2) + \cdots \\ &\quad + d_B(ga_1 \cdots a_{n-1}, ga_1 \cdots a_n) \\ &= d_B(e, a_1) + d_B(e, a_2) + \cdots + d_B(e, a_n) \\ &\leq c_1 \cdot n \\ &= c_1 \cdot d_A(g, h). \end{aligned}$$

De manera similar, intercambiando los papeles de A y B , se tiene que

$$c_2 = \max_{b \in B \cup B^{-1}} d_A(e, b)$$

es finito, distinto de cero, y satisface que

$$\frac{1}{c_2} \cdot d_A(g, h) \leq d_B(g, h).$$

Tomando $c = \max\{c_1, c_2\}$, se llega a que la identidad $Id : (G, d_A) \rightarrow (G, d_B)$ es una $(c, 0)$ -cuasi-isometría.

□

En otras palabras, el teorema anterior dice que en grupos finitamente generados, salvo cuasi-isometrías, solamente hay una métrica de las palabras. Un resultado inmediato del Teorema 3.8 es el siguiente.

Corolario 3.9. *Cualquier automorfismo de un grupo finitamente generado G es una cuasi-isometría.*

Demostración. Sea A un conjunto generador finito de G y sea $\Phi : G \rightarrow G$ un automorfismo. Dado que A es finito y Φ biyectivo, se cumple que $\Phi^{-1}(A)$ es un conjunto generador finito de G . Por el Teorema 3.8, la identidad es una cuasi-isometría entre (G, A) y $(G, \Phi^{-1}(A))$, por tanto existen constantes $c > 0$ y $b \geq 0$, tales que

$$\frac{1}{c} d_A(g, h) - b \leq d_{\Phi^{-1}(A)}(g, h) \leq c d_A(g, h) + b.$$

Por otro lado,

$$d_A(\Phi(g), \Phi(h)) = d_{\Phi^{-1}(A)}(g, h).$$

Sustituyendo esta igualdad en las últimas desigualdades se llega al resultado buscado.

□

3.2. Espacios Cuasi-Geodésicos

En la geometría de espacios métricos, a menudo resulta útil que el espacio en cuestión sea geodésico, es decir, que su métrica pueda ser realizada por trayectorias. Sin embargo, no siempre se tiene un espacio geodésico, es por esta razón que se introduce un concepto más general: el de un espacio cuasi-geodésico. Un espacio cuasi-geodésico es un espacio métrico que salvo un error uniforme, su métrica puede ser realizada por trayectorias.

3.2.1. Segmentos Geodésicos

Antes de abordar los conceptos de cuasi-geodésica y espacio cuasi-geodésico, es conveniente definir un segmento geodésico y un espacio geodésico.

Definición 3.10. Sean (X, d) un espacio métrico y $L \in \mathbb{R}_{\geq 0}$. Un **segmento geodésico** en X es un encaje isométrico $\gamma : [0, L] \rightarrow X$, donde el intervalo $[0, L]$ tiene la métrica inducida por la métrica usual de $\mathbb{R}$. La **longitud** de γ es L , y los puntos $\gamma(0)$ y $\gamma(L)$ se llaman **punto inicial** y **punto final** de γ , respectivamente.

Se dice que X es **geodésico** si para cualesquiera $x, x' \in X$ existe un segmento geodésico en X con punto inicial x y punto final x' .

El concepto de segmento geodésico definido anteriormente se relaciona con la noción de geodésica de la geometría diferencial (como la mencionada en el Capítulo 1), pero no son exactamente lo mismo. Los segmentos geodésicos en sentido métrico minimizan la longitud y distancia de forma global, mientras que en geometría diferencial únicamente se requiere que las geodésicas sean localmente isométricas y no globalmente.

Ejemplo 3.11. Para cada $n \in \mathbb{N}$ el espacio $\mathbb{R}^n$ con la métrica usual es geodésico. En efecto, si r es la distancia de cualesquiera dos puntos $x, y \in \mathbb{R}^n$ el segmento de recta

$x(r - t) - yt$ para $t \in [0, r]$ es un segmento geodésico entre x y y .

Ejemplo 3.12. El espacio $\mathbb{R}^2 - \{(0, 0)\}$ con la métrica inducida por la métrica Euclidiana de $\mathbb{R}^2$ *no* es geodésico. Por ejemplo, cualesquiera dos puntos antípodas no pueden ser unidos por un segmento de recta.

Ejemplo 3.13. La esfera $\mathbb{S}^2$ con la métrica Riemanniana estándar es un espacio geodésico. Los segmentos geodésicos son arcos de círculos máximos en $\mathbb{S}^2$. Obsérvese que los puntos antípodas en $\mathbb{S}^2$ pueden ser unidos por una infinidad de segmentos geodésicos distintos.

Ejemplo 3.14. El plano hiperbólico $\mathbb{H}$ es un espacio métrico geodésico. En el modelo del disco de Poincaré, los segmentos geodésicos son arcos de círculos que cortan ortogonalmente a $\mathbb{S}^1$.

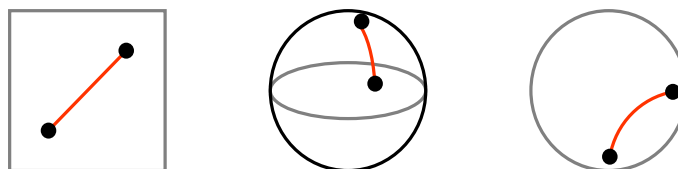


Figura 3.1: Espacios geodésicos y ejemplos de segmentos geodésicos.

3.2.2. Segmentos Cuasi-Geodésicos

Los grupos (no triviales) finitamente generados junto con la métrica de las palabras *no son* espacios geodésicos, pues el espacio métrico subyacente es discreto. Pero la magia de la geometría a gran escala permite que casi sean geodésicos.

Definición 3.15. Sea (X, d) un espacio métrico, y sean $c \in \mathbb{R}_{>0}$, y $b \in \mathbb{R}_{\geq 0}$. Una (c, b) -cuasi-geodésica en X es un encaje (c, b) -cuasi-isométrico $\gamma : I \rightarrow X$, donde I es algún intervalo cerrado $[t, t']$ de $\mathbb{R}$; el punto $\gamma(t)$ es el *punto inicial* de γ , y $\gamma(t')$ es el *punto final* de γ .

El espacio X es (c, b) -**cuasi-geodésico** si para todo $x, x' \in X$ existe una (c, b) -cuasi-geodésica en X con punto inicial en x y punto final x' .

Evidentemente cada espacio geodésico también es un espacio cuasi-geodésico (por ejemplo, $(1, 0)$ -cuasi-geodésico), pero en general, no todo espacio cuasi-geodésico es geodésico. A continuación se muestra el caso de un espacio cuasi-geodésico que no es geodésico.

Ejemplo 3.16. Para cada $\epsilon \in \mathbb{R}_{>0}$ el espacio $\mathbb{R}^2 - \{(0, 0)\}$ es $(1, \epsilon)$ -cuasi-geodésico (ver Figura 3.2.2) respecto a la métrica inducida por la métrica Euclidiana en $\mathbb{R}^2$, pero no es geodésico, puesto que cualquier encaje de un intervalo que una dos puntos antípodas tendría que pasar necesariamente por $(0, 0)$.

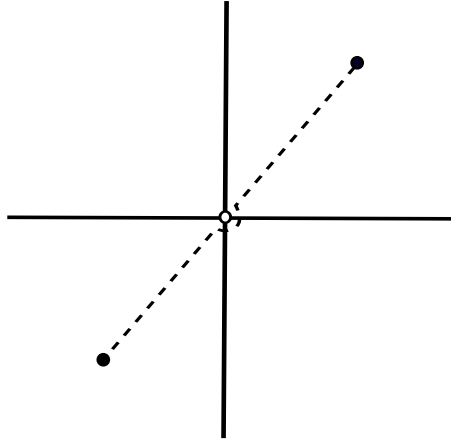


Figura 3.2: Una $(1, \epsilon)$ -cuasi-geodésica en $\mathbb{R}^2$

Ejemplo 3.17. Si G es un grupo y A es un conjunto generador de G , entonces (G, d_A) es un espacio $(1, 1)$ -cuasi-geodésico. En efecto, para cada $g, h \in G$ considérese $d_A(g, h) = n$. Así, $g^{-1}h = a_1 \dots a_n$, para ciertos $a_1, \dots, a_n \in A$. Luego, la función $\gamma : [1, n] \rightarrow G$ que a cada $t \in [1, n]$ le asigna $a_1 \dots a_{\lceil t \rceil}$, donde $\lceil t \rceil$ es el techo de t (es decir, el entero mayor más próximo a t). Haciendo algunas cuentas, se llega a que γ es una $(1, 1)$ -cuasi-geodésica.

3.3. Lema de Švarc-Milnor

El lema de Švarc-Milnor es uno de los resultados más poderosos de la teoría geométrica de grupos; burdamente dice que dada una acción “bonita” de un grupo sobre un espacio métrico “bonito”, se puede afirmar que el grupo es finitamente generado y que el grupo es cuasi-isométrico al espacio métrico en el que actúa.

En la práctica este resultado puede ser aplicado en ambos lados, es decir, si se quiere saber más acerca de la geometría de un grupo o saber si un grupo es finitamente generado, es suficiente exhibir una buena acción de este grupo sobre un espacio adecuado. O bien, siguiendo la filosofía del Programa de Erlangen, para conocer más de la geometría de este espacio es suficiente encontrar una buena acción de un grupo adecuado en el espacio métrico de interés. Por tal razón, el lema de Švarc-Milnor también es conocido como el “Lema Fundamental de la Teoría Geométrica de Grupos”.

La siguiente proposición es una versión métrica del lema de Švarc-Milnor, a partir de esta versión se deduce una más topológica, la cual es más utilizada en aplicaciones.

Lema 3.18 (Lema de Švarc-Milnor). *Sea G un grupo que actúa por isometrías sobre un espacio métrico (X, d_X) no vacío. Supóngase que X es un espacio (c, b) -cuasi-geodésico y que existe un subconjunto $B \subset X$ con las siguientes propiedades:*

- *El diámetro de B , denotado por $\text{diám } B$, es finito.*
- *Las G -traslaciones de B cubren a X , es decir*

$$\bigcup_{g \in G} g \cdot B = X.$$

- *El conjunto $S := \{g \in G \mid g \cdot B' \cap B' \neq \emptyset\}$ es finito, donde*

$$B' := B_{2,b}^{X,d_X}(B) = \{x \in X \mid \text{Existe } y \in B, \text{ tal que } d(x, y) \leq 2 \cdot b\}.$$

Entonces se cumplen las siguientes propiedades:

1. El grupo G es generado por S ; en particular, G es finitamente generado.
2. Para todo $x \in X$ la función asociada

$$\begin{aligned} G &\longrightarrow X \\ g &\longmapsto g \cdot x \end{aligned}$$

es una cuasi-isometría (respecto a la métrica d_S en G).

Demostración. Primero se demostrará que S es un conjunto generador de G (ver Definición 1.1).

Sea $g \in G$ un elemento arbitrario. Utilizado una cuasi-geodésica adecuada y siguiendo G -traslaciones de B a lo largo de esta cuasi-geodésica (ver Figura 3.3), se probará $g \in \langle S \rangle_G$.

Sea $x \in B$ fijo, dado que X es (c, b) -cuasi-geodésico, existe una (c, b) -cuasi-geodésica $\gamma : [0, L] \rightarrow X$ con punto inicial x y punto final $g \cdot x$.

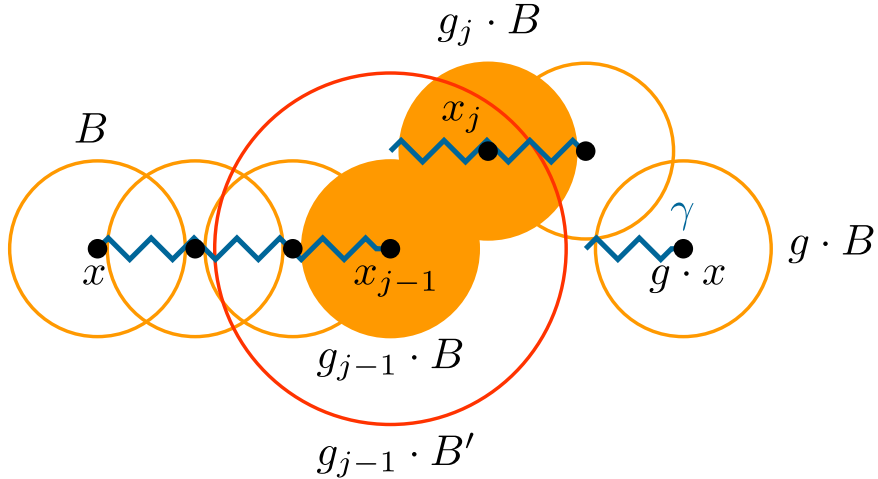


Figura 3.3: Una cuasi-geodésica γ cubierta por G -traslaciones de B .

Lo siguiente es buscar puntos suficientemente cercanos sobre esta cuasi-geodésica. Para esto, tómesese $n = \min\{k \in \mathbb{N} \mid L \cdot c/b \leq k\}$. Luego, para cada $j \in \{0, 1, \dots, n-1\}$

defínanse

$$t_j := j \cdot \frac{b}{c}, \text{ y } t_n := L,$$

de tal modo que $t_j \in [0, L]$. Esto es realizable puesto que para $j = 0, 1, \dots, n-1$, se cumple que

$$j \cdot \frac{b}{c} \leq L.$$

Luego, para cada $j = 0, 1, \dots, n$, defínase

$$x_j := \gamma(t_j).$$

Con las consideraciones anteriores se tiene que $x_0 = \gamma(0) = x$, $x_n = \gamma(L) = g \cdot x$ y que $x_j \in \gamma([0, L])$, para cada $j \in \{0, \dots, n-1\}$.

Por otro lado, como las G -traslaciones de B cubren a X , existen elementos $g_j \in G$ tales que $x_j \in g_j \cdot B$, es decir, elementos g_j del grupo tales que la traslación de B por g_j contiene a x_j . En particular, se puede elegir $g_0 := e$ y $g_n := g$.

Ahora se probará que para todo $j \in \{1, \dots, n\}$, el elemento del grupo $s_j := g_{j-1}^{-1} \cdot g_j$ se encuentra en S . Dado que γ es una (c, b) -cuasi-geodésica, se satisface que

$$\frac{1}{c} \cdot |t_{j-1} - t_j| - b \leq d(x_{j-1}, x_j) \leq c \cdot |t_{j-1} - t_j| + b.$$

De la segunda desigualdad y puesto que $|t_{j-1} - t_j| = |(j-1)\frac{b}{c} - j\frac{b}{c}| = \frac{b}{c}$, se obtiene

$$d(x_{j-1}, x_j) \leq c \cdot |t_{j-1} - t_j| + b \leq c \cdot \frac{b}{c} + b = 2 \cdot b.$$

Dado que $x_{j-1} \in g_{j-1} \cdot B$, se sigue que

$$x_j \in B_{2b}^{X,d}(g_{j-1} \cdot B) = \{x \in X \mid \text{Existe } y \in g_{j-1} \cdot B, \text{ tal que } d(x, y) \leq 2 \cdot b\}.$$

Como G actúa por isometrías sobre G , se sigue que

$$B_{2,b}^{X,d}(g_{j-1} \cdot B) = g_{j-1} \cdot B_{2,b}^{X,d}(B) = g_{j-1} \cdot B'.$$

Por otro lado, $x_j \in g_j \cdot B \subset g_j \cdot B'$, y por tanto

$$g_{j-1} \cdot B' \cap g_j \cdot B' \neq \emptyset.$$

Actuando ahora por g_{j-1}^{-1} se obtiene que

$$g_{j-1}^{-1} \cdot g_j \cdot B' \cap B' \neq \emptyset.$$

De la definición de S , se deduce que $s_j = g_{j-1}^{-1} \cdot g_j \in S$. Por lo tanto, se tiene que

$$g = g_n = g_{n-1} \cdot g_{n-1}^{-1} \cdot g_n = \cdots = g_0 \cdot (g_0^{-1} \cdot g_1) \cdot (\cdots) \cdot (g_{n-1}^{-1} \cdot g_n) = g_0 \cdot s_1 \cdots s_n$$

es un elemento del subgrupo generado por S . Como g es un elemento arbitrario, se sigue que S es un conjunto generador de G . En particular G es finitamente generado, pues por definición S es finito.

Ahora se probará que G con la métrica de las palabras inducida por S es cuasi-isométrico a X .

Sea $x \in X$, se probará que la función

$$\varphi : G \longrightarrow X$$

$$g \longmapsto g \cdot x,$$

es una cuasi-isometría demostrando que es un encaje cuasi-isométrico con imagen

cuasi-densa (Proposición 3.5).

Primero se probará que φ tiene imagen cuasi-densa. Sea $x' \in X$, dado que las G -traslaciones de B cubren a X , existen $g, g' \in G$ tales que $x \in g \cdot B$ y $x' \in g' \cdot B$. Luego, $g'g^{-1} \cdot x \in g' \cdot B$, lo que implica que

$$\begin{aligned} d(x', \varphi(g'g^{-1})) &= d(x', g'g^{-1} \cdot x) \\ &\leq \text{diám } g' \cdot B. \end{aligned}$$

Puesto que G actúa por cuasi-isometrías sobre X , el diámetro de $g' \cdot B$ es igual al diámetro de B , el cual por definición es finito. Por tanto, φ tiene imagen cuasi-densa.

Finalmente se probará que φ es un encaje cuasi-isométrico. Sea $g \in G$, primero se exhibirá una cota inferior uniforme de $d(\varphi(e), \varphi(g))$ en términos de $d_S(e, g)$.

Se puede suponer que x está contenido en B , puesto que G actúa por isometrías en X y que las G -traslaciones de B cubren a X , así se está como en la primera parte de la demostración. Sea $\gamma : [0, L] \rightarrow X$ una (c, b) -cuasi-geodésica de x a $g \cdot x$. Por un lado se tiene que

$$d(\varphi(e), \varphi(g)) = d(x, g \cdot x) = d(\gamma(0), \gamma(L)).$$

Por otro lado, como γ es un encaje (c, b) -cuasi-isométrico, se cumple que

$$d(\gamma(0), \gamma(L)) \geq \frac{1}{c}L - b.$$

Tomando $n = \min\{k \in \mathbb{N} \mid L \cdot c/b \leq k\}$ (como en la primera parte de la demostración), se llega a que $L \geq b(n-1)/c$, lo que implica que

$$\frac{1}{c}L - b \geq \frac{1}{c} \frac{b(n-1)}{c} - b = \frac{b}{c^2}n - \frac{b}{c^2} - b$$

Ahora bien, cuando se probó que S es un conjunto generador de G , se mostró que g se puede escribir con n elementos de S , lo que implica que $d_S(e, g) \leq n$. Por tanto

$$d(\varphi(e), \varphi(g)) \geq \frac{b}{c^2} d_S(e, g) - \frac{b}{c^2} - b.$$

Antes de obtener una cota superior uniforme de $d(\varphi(e), \varphi(g))$ en términos de $d_S(e, g)$, se probará que si $s \in S$, entonces $d(x, s \cdot x) \leq 2(\text{diám } B + 2b)$ (ver Figura 3.4).

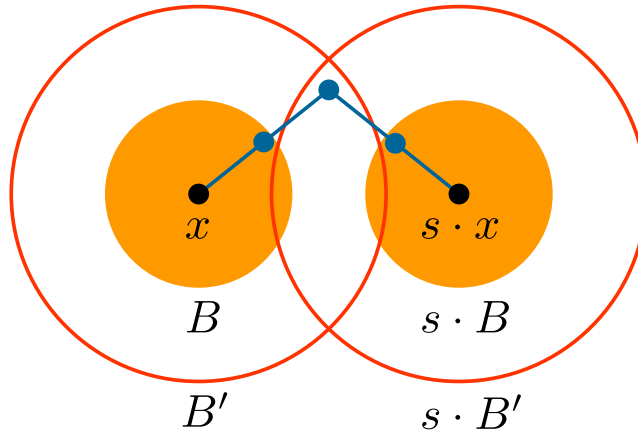


Figura 3.4: Si $s \in S$, entonces $d_s(x, s \cdot x) \leq 2 \cdot (\text{diám } B + 2b)$.

Si $y \in B'$, entonces existe $x' \in B$ tal que $d(x', y) \leq 2b$. Luego, por la desigualdad del triángulo $d(x, y) \leq d(x, x') + d(x', y) \leq \text{diám } B + 2b$. Puesto que la acción de G en X es por isometrías, de lo anterior se deduce que, para cualquier $y' \in s \cdot B'$, $d(y', s \cdot x) \leq \text{diám } B + 2b$. En particular, si $s \cdot B' \cap B' \neq \emptyset$, entonces existe $y \in s \cdot B' \cap B'$, tal que

$$d(x, s \cdot x) \leq d(x, y) + d(y, s \cdot x) \leq 2(\text{diám } B + 2b).$$

Sea $m = d_S(e, g)$, se sigue que existen $s_1, \dots, s_m \in S \cup S^{-1} = S$ tales que $g = s_1 \cdots s_m$. Usando la desigualdad del triángulo, el hecho de que G actúa por isometrías

en X , y que $s_j \cdot B' \cap B' \neq \emptyset$ para todo $j \in \{1, \dots, n-1\}$ se obtiene que

$$\begin{aligned}
d(\varphi(e), \varphi(g)) &= d(x, g \cdot x) \\
&\leq d(x, s_1 \cdot x) + d(s_1 \cdot x, s_1 \cdot s_2 \cdot x) + \dots \\
&\quad + d(s_1 \cdot \dots \cdot s_{m-1} \cdot x, s_1 \cdot \dots \cdot s_m \cdot x) \\
&= d(x, s_1 \cdot x) + d(x, s_2 \cdot x) + \dots + d(x, s_m \cdot x)m \\
&\leq 2(\text{diám } B + 2b)m \\
&= 2(\text{diám } B + 2b)d_S(e, g).
\end{aligned}$$

Luego, como el diámetro de B es finito y no depende de g , se sigue que $2(\text{diám } B + 2b)$ es finito y mayor que cero.

Después,

$$d(\varphi(g), \varphi(h)) = d(g \cdot x, h \cdot x) = d(x, g^{-1}h \cdot x) = d(\varphi(e), \varphi(g^{-1} \cdot h)) \quad \text{y} \quad d_S(g, h) = d_S(e, g^{-1} \cdot h)$$

se cumple para todo $g, h \in G$. Por tanto las cotas encontradas muestran que φ es un encaje cuasi-isométrico. $\square$

La versión (más) topológica del Lema de Švarc-Milnor se deduce como un corolario del Lema 3.18, pero antes de formularla es lícito recodar la definición de un espacio métrico *propio*.

Definición 3.19. Un espacio métrico (X, d) es **propio** si para todo $x \in X$ y todo $r \in \mathbb{R}_{>0}$ la bola cerrada $\{y \in X \mid d(x, y) \leq r\}$ es compacta respecto a la topología inducida por la métrica.

Corolario 3.20 (Versión topológica del Lema de Švarc-Milnor). *Sea G un grupo que actúa por isometrías sobre un espacio métrico (X, d) propio, geodésico y no vacío. Si esta acción es propiamente discontinua y cocompacta, entonces G es finitamente*

generado, y para todo $x \in X$ la función

$$\begin{aligned} G &\longrightarrow X \\ g &\longmapsto g \cdot x \end{aligned}$$

es una cuasi-isometría.

Demostración. El objetivo es utilizar el Lema de Švarc-Milnor en su versión anterior, para esto, primero nótese que como X es geodésico, en particular es un espacio $(1, b)$ -cuasi-geodésico para cualquier $b \in \mathbb{R}_{>0}$.

Ahora bien, dado que la acción de G es cocompacta, existe un compacto $B \subset X$ tal que las G -traslaciones de B cubren a X .

Además, el diámetro de B es finito, pues es un subconjunto compacto de un espacio métrico, y por tanto acotado. Por otro lado, $B' := B_{2 \cdot b}(B)$ también tiene diámetro finito. Como X es un espacio propio, se sigue que B' es compacto. En virtud de que la acción de G en X es propiamente discontinua se sigue que el conjunto $\{g \in G \mid g \cdot B' \cap B' \neq \emptyset\}$ es finito.

Así se satisfacen todas las hipótesis de la primera versión del Lema de Švarc-Milnor. Por tanto G es finitamente generado, y para todo $x \in X$ la función

$$\begin{aligned} G &\longrightarrow X \\ g &\longmapsto g \cdot x \end{aligned}$$

es una cuasi-isometría. □

Capítulo 4

Ligados en el Infinito

Cuando se define una relación en un grupo, naturalmente surge la pregunta de si tal relación se preserva bajo automorfismos. En esta sección se define la relación de *ligados en el infinito* en el grupo fundamental de una superficie hiperbólica, y se prueba usando cuasi-isometrías que esta relación se preserva bajo automorfismos.

4.1. Métricas en $\pi_1(S_g, p)$

Un resultado estándar de la teoría de cubrientes es que dada una función cubriente $\rho : \tilde{X} \rightarrow X$, si $\tilde{X}$ es simplemente conexo, entonces para cada $x \in X$ se puede definir una biyección entre el grupo fundamental $\pi_1(X, x)$ y la fibra $\rho^{-1}(x)$ (ver Teorema B.35). Tal biyección se construye de la siguiente forma: Se toma un punto $x_0 \in \tilde{X}$ tal que $\rho(x_0) = x$, es decir, $x_0 \in \rho^{-1}(x)$. Luego, a cada $[\alpha] \in \pi_1(X, x)$, se le asocia el punto final de los levantamientos de α a $\tilde{X}$ con punto inicial en x_0 . Como $\tilde{X}$ es simplemente conexo, esta asignación resulta ser biyectiva (ver Teorema B.35). En particular, al elemento neutro e le corresponde el punto x_0 .

En el capítulo precedente se mencionó que todo conjunto generador de un grupo induce una métrica en tal grupo (la métrica de las palabras). En el caso del grupo

fundamental $\pi_1(S_g, p)$ de una superficie S_g de género al menos 2, la biyección entre $\pi_1(S_g, p)$ y la fibra de p permite inducir otra métrica, la *métrica hiperbólica*.

Definición 4.1 (Métrica hiperbólica). Sean $g \geq 2$, $\rho : \mathbb{H} \rightarrow S_g$ una función cubriente y $p \in S_g$. La distancia entre dos elementos del $\pi_1(S_g, p)$ se define como la distancia hiperbólica entre sus correspondientes en la fibra $\rho^{-1}(p)$ (ver Figura 4.1).

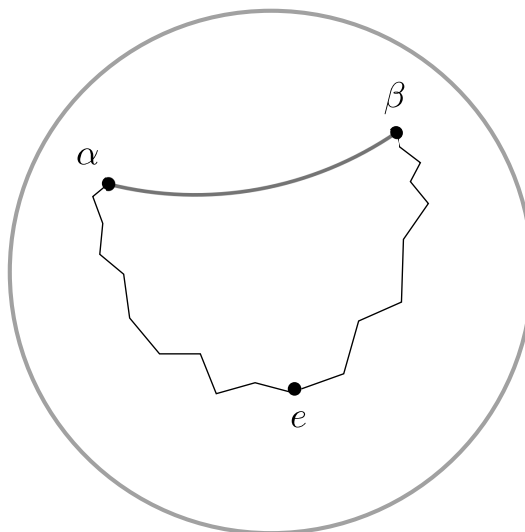


Figura 4.1: Levantamientos de $\alpha, \beta \in \pi_1(S_g, p)$ y la distancia hiperbólica entre ellos.

Nótese que la métrica hiperbólica en $\pi_1(S_g, p)$ depende directamente de la función cubriente ρ , en efecto, si se tiene otra función cubriente $\rho' : \mathbb{H} \rightarrow S_g$ los conjuntos $\rho^{-1}(p)$ y $\rho'^{-1}(p)$ no necesariamente son iguales.

Por otro lado, $\rho^{-1}(p)$ es un espacio métrico (con la métrica inducida por $\mathbb{H}$) que naturalmente es isométrico a $\pi_1(S_g, p)$ con la métrica hiperbólica inducida por ρ . Más aún, se puede definir una acción por isometrías del $\pi_1(S_g, p)$ en $\rho^{-1}(p)$: Para cada $[\alpha] \in \pi_1(S_g, p)$ y cada $q \in \rho^{-1}(p)$, $[\alpha] \cdot q$ es el punto final de los levantamientos de α con punto inicial en q . Se puede probar que tal acción es por isometrías, libre y propiamente discontinua. Aplicando el lema de Švarc-Milnor (Corolario 3.20), se deduce que $\pi_1(S_g, p)$ con la métrica de palabras es cuasi-isométrico a $\rho^{-1}(p)$, y por

tanto cuasi-isométrico a $\pi_1(S_g, p)$ con la métrica hiperbólica.

Ahora bien, el $\pi_1(S_g, p)$ con cada métrica hiperbólica es cuasi-isométrico al $\pi_1(S_g, p)$ con alguna métrica hiperbólica. Por el Teorema 3.8, salvo cuasi-isometrías, solamente hay una métrica de las palabras en $\pi_1(S_g, p)$. Por tanto estás dos métricas en $\pi_1(S_g, p)$ que salen de forma natural son equivalentes bajo cuasi-isometría. Este argumento es muy útil para cambiar entre la métrica de palabras y la métrica hiperbólica en $\pi_1(S_g, p)$. Por ejemplo, el Corolario 3.9 se demostró usando la métrica de las palabras, y será aplicado posteriormente en la demostración del Lema 4.5 usando la métrica hiperbólica.

4.2. Ligados en el Infinito

En la Sección 1.2.2 se vio que una superficie es hiperbólica si admite al plano hiperbólico como su cubriente universal (riemanniano), lo cual siempre sucede para una superficie S_g de género al menos 2.

Si $(\mathbb{H}, \rho)$ es un espacio cubriente fijo de S_g , una **transformación de cubierta** es un homeomorfismo $\eta : \mathbb{H} \rightarrow \mathbb{H}$, tal que $\rho \circ \eta = \rho$. La acción del $\pi_1(S_g, p)$ sobre $\rho^{-1}(p)$, mencionada en la sección anterior, se puede extender continuamente al todo el plano hiperbólico $\mathbb{H}$, de tal forma que $\pi_1(S_g, p)$ actúa por isometrías en $\mathbb{H}$ mediante transformaciones de cubierta, de forma libre y propiamente discontinua. (Para más detalles véase el Apéndice B). En particular $\pi_1(S_g, p)$ está encajado en el grupo de isometrías de $\mathbb{H}$, pues su acción en $\mathbb{H}$ es libre. En otras palabras, a cada elemento del $\pi_1(S_g, p)$ le corresponde una única isometría de $\mathbb{H}$ y si a dos elementos les corresponde la misma isometría entonces son iguales.

Definición 4.2. Sea S es una superficie hiperbólica, se dice que un elemento $\gamma \in \pi_1(S, p)$ es **hiperbólico** si la transformación de cubierta correspondiente a γ es una isometría hiperbólica.

Una isometría hiperbólica de $\mathbb{H}$ se caracteriza por dejar únicamente dos puntos fijos en la frontera de $\mathbb{H}$. Puesto que las isometrías mandan geodésicas en geodésicas, se sigue que una isometría hiperbólica deja fija (como conjunto) a una única geodésica; a tal geodésica se le llama el **eje** de la isometría. Uno de los extremos del eje de una isometría se llama **atractor** y el otro **repulsor** (véase la Figura 4.2).

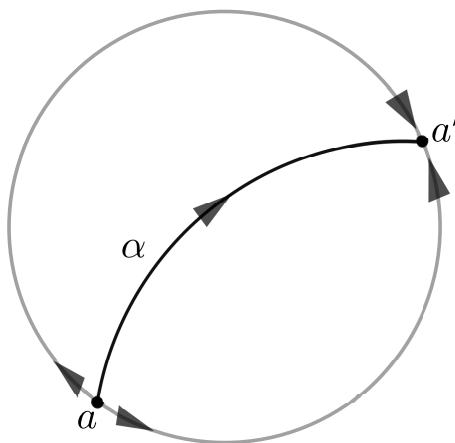


Figura 4.2: a es el punto atractor y a' el punto repulsor de eje de α .

Definición 4.3. Sean S una superficie hiperbólica y $p \in S$. El **eje** de un elemento hiperbólico de $\pi_1(S, p)$ es el eje de su correspondiente isometría hiperbólica.

Definición 4.4. Se dice que dos elementos $\gamma, \delta \in \pi_1(S_g, p)$ están **ligados en el infinito** si sus ejes se cortan en el interior de $\mathbb{H}$ (ver Figura 4.3).

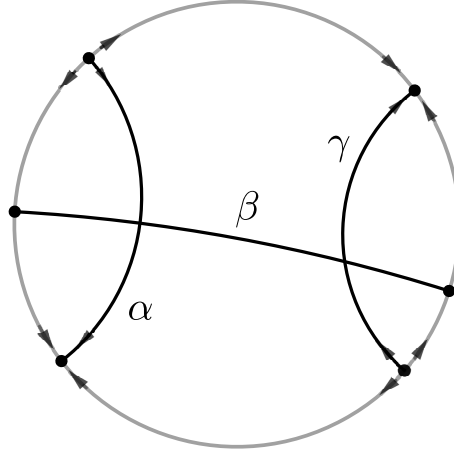


Figura 4.3: β está ligado en el infinito con α y con γ , pero α y γ no están ligados en el infinito.

Lema 4.5. Sea $g \geq 2$ y sea $\rho : \mathbb{H} \rightarrow S_g$ una función cubriente fija. Sea Φ un automorfismo de $\pi_1(S_g, p)$ y sean γ y δ elementos no triviales de $\pi_1(S_g, p)$. Entonces los elementos $\Phi(\gamma)$ y $\Phi(\delta)$ están ligados en el infinito si y sólo si γ y δ están ligados en el infinito.

Demostración. Dado que S_g es una superficie hiperbólica cerrada, todos los elementos (no triviales) del grupo fundamental de S_g son hiperbólicos, y por lo tanto tiene sentido hablar sobre los elementos hiperbólicos del grupo fundamental que están ligados en el infinito. Obsérvese que como Φ es invertible, es suficiente probar que si γ y δ son dos elementos que *no están* ligados en el infinito, entonces $\Phi(\gamma)$ y $\Phi(\delta)$ *no están* ligados en el infinito. Además, se puede suponer que los ejes γ y δ no son el mismo, pues esto es equivalente a que uno es una potencia (no trivial) del otro, y esta propiedad se preserva por automorfismos.

Por el Corolario 3.9 y el Corolario 3.20, Φ es una cuasi-isometría de $\pi_1(S_g, p)$. Supóngase que Φ es una (K, C) -cuasi-isometría, respecto a la métrica hiperbólica inducida por la función cubriente ρ , con $K \geq 1$ y $C \geq 0$. Sea D el diámetro de algún dominio fundamental fijo para $\pi_1(S_g, p)$ en $\mathbb{H}$. Nótese que D es finito, pues es el diámetro de un conjunto compacto, dado que la acción es cocompacta.

Considérese una constante $R > 2DK^2 + 2CK$, un punto fijo x_0 en $\rho^{-1}(p)$, y la órbita de x_0 respecto de γ :

$$\mathcal{O}_\gamma = \{\gamma^k \cdot x_0 \mid k \in \mathbb{Z}\}.$$

Conéctese los puntos de $\mathcal{O}_\gamma$ por un camino geodésico a trozos que se esté completamente dentro de una vecindad fija del eje correspondiente a γ , de tal forma que cada segmento geodésico del camino una dos puntos en la órbita de x_0 que están en dominios fundamentales adyacentes. Denótese a tal camino por su conjunto de vértices $\{\alpha_i\}$.

Dado que γ y δ son isometrías hiperbólicas de $\mathbb{H}$ que no están ligados en el infinito, y por construcción el camino $\{\alpha_i\}$ se encuentra dentro de una vecindad fija del eje de γ , se puede encontrar una constante $N = N(R)$ tal que cada punto de

$$\mathcal{O}_{\delta^N} = \{(\delta^N)^k \cdot x_0 \mid k \in \mathbb{Z}, \text{ con } k \neq 0\}$$

está a una distancia de al menos $R + D$ de cada punto de $\mathcal{O}_\gamma$. Nótese que $\mathcal{O}_{\delta^N}$ no es la órbita completa de x_0 respecto a δ^N puesto que no contiene al punto x_0 . Análogamente, conéctese los puntos de $\mathcal{O}_{\gamma^N}$ por caminos $\{\beta_i\}$ donde cada β_i está en la órbita de x_0 , tal que el camino β_i está a una distancia hiperbólica mayor o igual a R de los caminos $\{\alpha_i\}$ y de manera que los vértices consecutivos β_i y β_{i+1} pertenezcan a dominios fundamentales adyacentes. Para encontrar el camino $\{\beta_i\}$, se puede iniciar con cualquier camino continuo bi-infinito que conecte los puntos de $\mathcal{O}_{\delta^N}$ y se mantenga fuera de la $(R + D)$ -vecindad del camino $\{\alpha_i\}$, finalmente se unen los puntos de la órbita de x_0 que están en dominios fundamentales adyacentes por los que pasa este camino.

Obsérvese que en los dos caminos contruidos, la longitud de cada segmento

geodésico es a lo más $2D$, debido a que para cualesquiera dos puntos en dominios fundamentales adyacentes su distancia es menor o igual a $2D$. Además, como los vértices de cada camino se encuentran en dominios fundamentales distintos, estos se pueden identificar con elementos particulares de $\pi_1(S_g, p)$, es decir, cada α_i y β_i representan (o son) elementos distintos de $\pi_1(S_g, p)$, y por el mismo argumento $\Phi(\alpha_i)$ y $\Phi(\beta_i)$ representan (o son) puntos distintos en la órbita de x_0 .

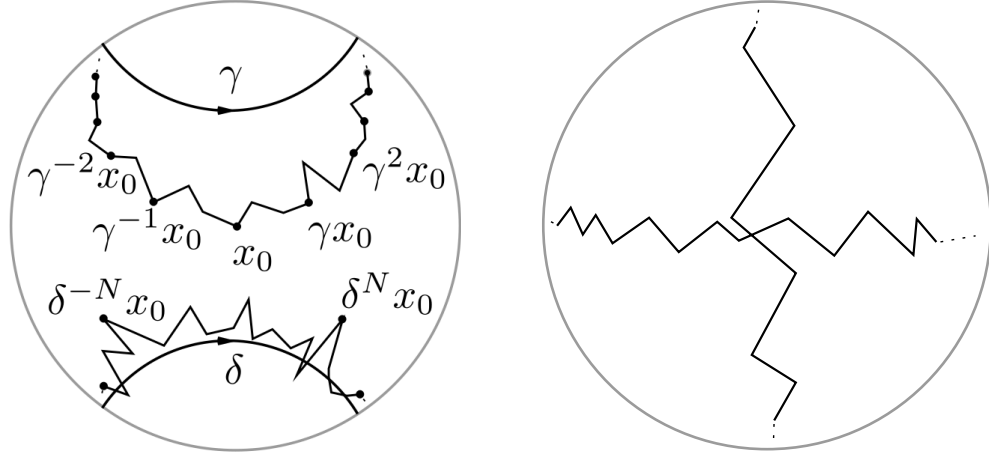


Figura 4.4: En la figura de la izquierda se observan los caminos poligonales construidos en la demostración del Lema 4.5; y en el de derecha, caminos poligonales que están ligados en el infinito.

Prosiguiendo por contradicción, supóngase que los elementos $\Phi(\gamma)$ y $\Phi(\delta)$ están ligados en el infinito. Dado que Φ es una (K, C) -cuasi-isometría, cada segmento cuasi-geodésico de $\{\Phi(\alpha_i)\}$ y $\{\Phi(\beta_i)\}$ tiene longitud a lo más $K(2D) + C$: Por la suposición de que estos caminos se cortan, dos de los segmentos geodésicos de cada camino se deben cortar (ver Figura 4.4). Ahora bien, cada uno de estos segmentos geodésicos tiene al menos un punto final cuya distancia al punto de intersección es a lo más $(K(2D) + C)/2$, por la desigualdad de triángulo, se sigue que estos puntos están a una distancia menor o igual $K(2D) + C$.

Por la biyección entre la órbita de x_0 y el $\pi_1(S_g, p)$, existen dos elementos $\gamma', \delta' \in \pi_1(S_g, p)$ tales que $d(\gamma', \delta') \geq R$ y $d(\Phi(\gamma'), \Phi(\delta')) \leq K(2D) + C$. Por otro lado, se

tiene que $R > 2DK^2 + 2CK$, de manera que

$$d(\gamma', \delta') \geq R > 2DK^2 + 2CK = K(K(2D) + C + C) \geq K(d(\Phi(\gamma'), \Phi(\delta')) + C),$$

lo que implica que

$$\frac{1}{K} \cdot d(\gamma', \delta') - C > d(\Phi(\gamma'), \Phi(\delta'));$$

esto contradice a la suposición de que Φ es una cuasi-isometría con constantes K y C . Por lo tanto, $\Phi(\gamma)$ y $\Phi(\delta)$ no pueden estar ligados en el infinito, que es lo que se quería demostrar. $\square$

Además de elementos ligados en el infinito, también se puede hablar de cuando dos elementos hiperbólicos $\alpha, \beta \in \pi_1(S_g, p)$ se encuentran en el *mismo lado* de un elemento hiperbólico $\gamma \in \pi_1(S_g, p)$.

Definición 4.6. Sean α, β y γ elementos hiperbólicos de $\pi_1(S, g)$. Supóngase que α y β no están ligados en el infinito (y no comparten un extremo de sus ejes) con γ , se dice que α y β están del **mismo lado** de γ si los ejes correspondientes a α y a β se encuentran completamente en un lado del eje de γ , de lo contrario se dice que están en **lados opuestos**.

Corolario 4.7. Sea $g \geq 2$ y $\mathbb{H} \rightarrow S_g$ una función cubriente fija. Sea Φ un automorfismo de $\pi_1(S_g, p)$. Si α, β y γ son elementos de $\pi_1(S_g, p)$ con ejes distintos, entonces los ejes de $\Phi(\alpha)$ y $\Phi(\beta)$ están del mismo lado de $\Phi(\gamma)$ si y sólo si α y β están el mismo lado de γ .

Demostración. Nótese que los ejes de correspondientes a α y β se encuentran en el mismo lado del eje de γ si y sólo si existe un elemento $\delta \in \pi_1(S_g, p)$ que está ligado en el infinito con α y β pero no con γ . Aplicando el Lema 4.5 queda demostrado este resultado. $\square$

Capítulo 5

Teorema de Dehn-Nielsen-Baer

El grupo modular de Teichmüller fue introducido por primera vez en 1920 por el matemático alemán Max Dehn [10]. Este es relativamente un nuevo campo de estudio en donde la topología y el álgebra se mezclan perfectamente produciendo resultados bastante interesantes.

De entre todos los resultados, uno de especial interés es el *teorema de Dehn-Nielsen-Baer*, el cual declara que el grupo modular de Teichmüller de una superficie S_g , denotado por $\text{Mod}(S_g)$, es isomorfo a un subgrupo de índice dos del grupo de automorfismos externos de $\pi_1(S_g, p)$. La prueba original de tal teorema se atribuye a Dehn, aunque fue Nielsen el primero en publicar una demostración en 1927. Baer fue el primero en demostrar que existe una inyección del $\text{Mod}(S_g)$ en el $\text{Out}(S_g)$.

5.1. Teorema de Dehn-Nielsen-Baer

5.1.1. Grupo Modular de Teichmüller Extendido

Definición 5.1. Sea S una superficie sin frontera. El **grupo modular (de Teichmüller) extendido**, denotado por $\text{Mod}^\pm(S)$, es el grupo de *clases de isotopía* de todos los homeomorfismos de S , incluyendo los que no preservan la orientación.

Un resultado inmediato es que el grupo $\text{Mod}(S)$ es un subgrupo de índice dos del grupo $\text{Mod}^\pm(S)$. Se puede definir un homomorfismo de grupos sobreyectivo de $\psi : \text{Mod}^\pm(S) \rightarrow \mathbb{Z}/2\mathbb{Z}$ tal que a cada clase de isotopía se le asigna 0 si preserva la orientación, o 1 si no la preserva. Es fácil ver que el kernel de ψ es $\text{Mod}(S)$, por el primer teorema de isomorfismos se deduce que $\text{Mod}^\pm(S)/\text{Mod}(S) \cong \mathbb{Z}/2\mathbb{Z}$. De lo anterior se sigue que el índice del $\text{Mod}(S)$ en $\text{Mod}^\pm(S)$ es dos. Con este resultado se deducen inmediatamente los siguientes ejemplos:

Ejemplo 5.2. $\text{Mod}^\pm(S^2) \cong \mathbb{Z}/2\mathbb{Z}$.

Ejemplo 5.3. Puesto que el $\text{Mod}(T^2) \cong \text{SL}(2, \mathbb{Z})$, se tiene que

$$\text{Mod}^\pm(T^2) \cong \text{GL}(2, \mathbb{Z}).$$

5.1.2. Grupo de Automorfismos Externos

Definición 5.4. Sea G un grupo y $\text{Aut}(G)$ el grupo de automorfismos de G . Para cada $h \in G$, se llama **automorfismo interno** de G al automorfismo $I_h : G \rightarrow G$ dado por $g \mapsto h \cdot g \cdot h^{-1}$ para todo $g \in G$. El **grupo de automorfismos internos** de G , denotado por $\text{Inn}(G)$, es el subgrupo normal de $\text{Aut}(G)$ que consiste de todos los automorfismos internos de G . El **grupo de automorfismos externos** de G es el cociente

$$\text{Out}(G) := \frac{\text{Aut}(G)}{\text{Inn}(G)}.$$

En otras palabras, $\text{Out}(G)$ es el grupo de automorfismos de G considerado hasta conjugación. Es importante notar que mientras un elemento del $\text{Out}(G)$ no actúa sobre los elementos de G , si actúa sobre el conjunto de las *clases de conjugación* de los elementos de G .

Nótese que $g^{-1} \circ f \in \text{Inn}(G)$ si y sólo si $f = I_a \circ g$ para algún $a \in G$. Así para probar que dos automorfismos f y g de G son representantes del mismo elemento en

el $\text{Out}(G)$, es suficiente encontrar un elemento $a \in G$ tal que $f = I_a \circ g$.

5.1.3. Homomorfismo Natural

Dada una superficie S y un punto $p \in S$, se busca establecer una relación entre el $\text{Mod}^\pm(S)$ y el $\text{Out}(\pi_1(S, p))$. Encontrar una relación es realmente fácil, puesto que las mismas definiciones propician un homomorfismo entre ambos grupos.

Para construir tal homomorfismo, primero se probará que toda función continua $\phi : S \rightarrow S$ induce un endomorfismo en el $\pi_1(S, p)$, y que todo homeomorfismo induce un automorfismo.

Proposición 5.5. *Sea S un espacio arco-conexo y sea un punto $p \in S$. Si $\phi : S \rightarrow S$ es una función continua y γ una curva que une a p con $\phi(p)$, entonces la función*

$$\phi_* : \pi_1(S, p) \rightarrow \pi_1(S, p)$$

*dada por $[\alpha] \mapsto [\gamma * \phi(\alpha) * \gamma^{-1}]$ es un endomorfismo.*

Demostración. Dado que ϕ es una función continua, para cualquier curva cerrada α en S se tiene que $\phi(\alpha) = \phi \circ \alpha$ también es una curva cerrada en S . Más aún, si α y β son dos curvas en S homotópicamente equivalentes (ver Apéndice B), entonces $\phi(\alpha)$ y $\phi(\beta)$ son homotópicamente equivalentes.

Por otro lado, dado que α está basada en p , se tiene que $\phi(\alpha)$ está basada en $\phi(p) = \gamma(1)$. Por tanto la curva $\gamma * \phi(\alpha) * \gamma^{-1}$ está basada en p y su clase de homotopía es un elemento de $\pi_1(S, p)$.

Lo mencionado anteriormente demuestra que ϕ_* es una función bien definida. Ahora se probará que ϕ_* también es un homomorfismo de grupos.

Considérese un par de curvas α y β basadas en p . Nótese que ϕ separa el producto

de curvas, es decir, $\phi(\alpha * \beta) = \phi(\alpha) * \phi(\beta)$. Así, se tiene lo siguiente.

$$\begin{aligned}
\phi_*([\alpha][\beta]) &= \phi_*([\alpha * \beta]) \\
&= [\gamma * \phi(\alpha * \beta) * \gamma^{-1}] \\
&= [\gamma * \phi(\alpha) * \phi(\beta) * \gamma^{-1}] \\
&= [\gamma * \phi(\alpha) * \gamma^{-1} * \gamma * \phi(\beta) * \gamma^{-1}] \\
&= [\gamma * \phi(\alpha) * \gamma^{-1}][\gamma * \phi(\beta) * \gamma^{-1}] \\
&= \phi_*([\alpha])\phi_*([\beta]).
\end{aligned}$$

Por lo tanto ϕ_* es un endomorfismo de $\pi_1(S, p)$. □

Se llamará a ϕ_* un **endomorfismo inducido** por ϕ .

Corolario 5.6. *Si $\phi : S \rightarrow S$ es un homeomorfismo, entonces cualquier endomorfismo inducido por ϕ es un automorfismo.*

Demostración. Considérese un camino γ que une a p con $\phi(p)$. Dado que ϕ es biyectiva y con inversa continua, $\phi^{-1}(\gamma)$ es un camino que une a $\phi^{-1}(p)$ con p . Defínase un endomorfismo $\phi_*^{-1} : \pi_1(S, p) \rightarrow \pi_1(S, p)$ inducido por ϕ^{-1} , mediante

$$[\beta] \mapsto [\phi^{-1}(\gamma^{-1}) * \phi^{-1}(\beta) * \phi^{-1}(\gamma)].$$

Recordando que ϕ separa el producto de curvas y realizando un cálculo fácil se llega a que $\phi_* \circ \phi_*^{-1} = \phi_*^{-1} \circ \phi_* = Id$. Esto demuestra que ϕ_*^{-1} es inverso ϕ_* , por lo tanto el homomorfismo inducido ϕ_* es un automorfismo de $\pi_1(S, p)$. Puesto que γ se tomó arbitrario, se sigue que cualquier endomorfismo inducido ϕ_* es un automorfismo de $\pi_1(S, p)$. □

Ahora bien, si $\phi : S \rightarrow S$ es un homeomorfismo fijo, las diferentes elecciones de γ inducen diferentes automorfismos ϕ_* de $\pi_1(S, p)$. La siguiente proposición demuestra

que estos automorfismos están relacionados por conjugación.

Proposición 5.7. *Sean S una superficie, $\phi : S \rightarrow S$ una función continua, y p un punto en S . Si $\phi_*^1, \phi_*^2 : \pi(S, p) \rightarrow \pi(S, p)$ son dos endomorfismos inducidos por ϕ , entonces para cualquier $[\alpha] \in \pi_1(S, p)$, $\phi_*^1([\alpha])$ y $\phi_*^2([\alpha])$ son conjugados en $\pi(S, p)$.*

Demostración. Sean γ_1 y γ_2 caminos arbitrarios que unen a p con $\phi(p)$. Sin pérdida de generalidad, se puede suponer que

$$\phi_*^1([\alpha]) = [\gamma_1 * \phi(\alpha) * \gamma_1^{-1}] \quad \text{y} \quad \phi_*^2([\alpha]) = [\gamma_2 * \phi(\alpha) * \gamma_2^{-1}],$$

para todo $[\alpha] \in \pi_1(S, p)$.

Nótese que $\gamma_1 * \gamma_2^{-1}$ es una curva cerrada basada en p , y que

$$\phi_*^1([\alpha]) = [\gamma_1 * \gamma_2^{-1}] \phi_*^2([\alpha]) [\gamma_2 * \gamma_1^{-1}].$$

Esto significa que para cada $[\alpha] \in \pi_1(S, p)$, $\phi_*^1([\alpha])$ y $\phi_*^2([\alpha])$ son conjugados en $\pi_1(S, p)$. □

Proposición 5.8. *Si $\phi, \varphi : S \rightarrow S$ son dos homeomorfismos isotópicos de S , entonces $\varphi_*^{-1} \circ \phi_*$ es un automorfismo interno de $\pi_1(S, p)$, para cualesquiera automorfismos inducidos ϕ_* y φ_* de $\pi_1(S, p)$.*

Demostración. Considérense γ_ϕ y γ_φ dos caminos que unen a p con $\phi(p)$ y $\varphi(p)$, respectivamente. Sea $F : I \times S \rightarrow S$ una isotopía entre ϕ y φ , y llámese γ al camino de $\phi(p)$ a $\varphi(p)$ dado por $\gamma(t) = F(t, p)$. Se tiene lo siguiente:

$$\begin{aligned} \phi_*([\alpha]) &= [\gamma_\phi * \gamma * \gamma_\varphi^{-1}] [\gamma_\varphi * \varphi(\alpha) * \gamma_\varphi^{-1}] [\gamma_\varphi * \gamma^{-1} * \gamma_\phi^{-1}] \\ &= [\gamma_\phi * \gamma * \gamma_\varphi^{-1}] \varphi_*([\alpha]) [\gamma_\varphi * \gamma^{-1} * \gamma_\phi^{-1}]. \end{aligned}$$

Puesto que $[\gamma_\phi * \gamma * \gamma_\phi^{-1}]$ es un elemento de $\pi_1(S, p)$, se sigue que $\varphi_*^{-1} \circ \phi_* \in \text{Inn}(\pi_1(S, p))$.

□

Una consecuencia inmediata de la Proposición 5.7 y el Corolario 5.6, es que si $\phi : S \rightarrow S$ es un homeomorfismo, entonces los automorfismos inducidos pertenecen a la misma clase lateral (izquierda) del $\text{Inn}(\pi_1(S, p))$, es decir, son representantes del mismo elemento en el $\text{Out}(\pi_1(S, p))$. En virtud de lo anterior y la Proposición 5.8, naturalmente se puede definir una función

$$\sigma : \text{Mod}^\pm(S) \rightarrow \text{Out}(\pi_1(S, p)),$$

con la asignación $[\phi] \mapsto [\phi_*]$.

Proposición 5.9. *La función σ es un homomorfismo de grupos.*

Demostración. Para demostrar que σ es un homomorfismo de grupos es suficiente probar que la composición de dos automorfismos inducidos es igual a un automorfismo inducido de la composición, es decir, $\varphi_* \circ \phi_* = (\varphi \circ \phi)_*$ para cualquier par de homeomorfismos φ y ϕ de S .

Considérese un camino γ_1 que une a p con $\varphi(p)$ y otro camino γ_2 que une a p con $\phi(p)$. Entonces, los automorfismos siguientes son automorfismos inducidos por φ y ϕ :

$$\varphi_*([\alpha]) = [\gamma_1 * \varphi(\alpha) * \gamma_1^{-1}] \quad \text{y} \quad \phi_*([\alpha]) = [\gamma_2 * \phi(\alpha) * \gamma_2^{-1}],$$

para todo $[\alpha] \in \pi_1(S, p)$. Se tiene lo siguiente.

$$\begin{aligned} \varphi_* \circ \phi_*([\alpha]) &= \varphi_*([\gamma_2 * \phi(\alpha) * \gamma_2^{-1}]) \\ &= [\gamma_1 * \varphi(\gamma_2 * \phi(\alpha) * \gamma_2^{-1}) * \gamma_1^{-1}] \\ &= [\gamma_1 * \varphi(\gamma_2) * \varphi(\phi(\alpha)) * \varphi(\gamma_2^{-1}) * \gamma_1^{-1}] \end{aligned}$$

Nótese que $\gamma_1 * \varphi(\gamma_2)$ es un camino que une a p con $\varphi \circ \phi(p)$, y que $[\varphi(\gamma_2)^{-1}] = [\varphi(\gamma_2^{-1})]$. Por tanto $(\varphi \circ \phi)_*([\alpha]) = [\gamma_1 * \varphi(\gamma_2) * \varphi(\phi(\alpha)) * \varphi(\gamma_2^{-1}) * \gamma_1^{-1}]$ es un automorfismo inducido por $\varphi \circ \phi$. En consecuencia de lo anterior se tiene que

$$\sigma([\phi][\varphi]) = \sigma([\phi \circ \varphi]) = [(\phi \circ \varphi)_*] = [\phi_* \circ \varphi_*] = [\phi_*][\varphi_*] = \sigma([\phi])\sigma([\varphi]).$$

Lo que demuestra que σ es un homomorfismo de grupos. □

5.1.4. Teorema de Dehn-Nielsen-Baer

Teorema 5.10 (Teorema de Dehn-Nielsen-Baer). *Sea S_g una superficie orientable de género $g \geq 2$. El homomorfismo de grupos*

$$\begin{aligned} \sigma : \text{Mod}^\pm(S_g) &\longrightarrow \text{Out}(\pi_1(S_g, p)), \\ [\phi] &\longmapsto [\phi_*], \end{aligned}$$

es un isomorfismo.

En la Proposición 5.9 se demostró que σ es un homomorfismo de grupos, así para demostrar el teorema de Dehn-Nielsen-Baer es suficiente demostrar que σ es un homomorfismo inyectivo y sobreyectivo. Primeramente se probará que σ es inyectivo y posteriormente que es sobreyectivo.

Nótese que si $\phi : S_g \rightarrow S_g$ es un homeomorfismo tal que $\sigma([\phi])$ es igual al elemento neutro de $\text{Out}(\pi_1(S_g, x_0))$, entonces cualquier automorfismo inducido ϕ_* es un automorfismo interno de $\pi_1(S_g, p)$, y por tanto ϕ_* fija todas las clases de conjugación

de $\pi_1(S_g, x_0)$. De la correspondencia biyectiva

$$\left\{ \begin{array}{c} \text{Clases no triviales} \\ \text{de conjugación} \\ \text{en } \pi_1(S, p) \end{array} \right\} \longleftrightarrow \left\{ \begin{array}{c} \text{Clases no triviales de} \\ \text{homotopía libre de curvas cerradas} \\ \text{y orientadas en } S \end{array} \right\}$$

se deduce que ϕ deja fijo cualquier clase de homotopía de curvas cerradas y orientadas en S . En particular, ϕ fija las clases de isotopía libre de cualquier conjunto de curvas que satisfacen las hipótesis del método de Alexander (Teorema 2.14) en consecuencia ϕ es isotópico a la función identidad en S_g . De esta manera, la clase de isotopía de ϕ es el elemento neutro del $\text{Mod}^\pm(S_g)$, de donde se sigue que el $\text{Ker}(\sigma)$ es trivial y por tanto σ es inyectiva.

Para finalizar con la demostración del teorema de Dehn-Nielsen-Baer solo falta probar que σ también es sobreyectiva.

Con la finalidad de facilitar la lectura, en lo consiguiente a cualquier clase de homotopía (o isotopía) $[\alpha]$ se denotará únicamente por α .

Demostración. Considérese $[\Phi]$ un elemento arbitrario de $\text{Out}(\pi_1(S_g, p))$ y sea Φ un automorfismo representante. Además, déjese fija una función cubriente $\rho : \mathbb{H} \rightarrow S_g$.

Sea $(c_1, c_2, \dots, c_{2g})$ una cadena de clases de isotopía de curvas cerradas simples en S . Como en la Subsección 1.3.3, esto significa que $i(c_i, c_{i+1}) = 1$ y $i(c_i, c_j) = 0$ si $|i - j| > 1$. Supóngase que cada c_i está orientada de tal forma que $\hat{i}(c_i, c_{i+1}) = +1$.

Recuérdese que clases de homotopía libre de curvas orientadas en S_g corresponden a clases de conjugación de elementos de $\pi_1(S_g, p)$. No es difícil ver que cada c_i corresponde a la clase de conjugación de un $\gamma_i \in \pi_1(S_g, p)$, donde $\{\gamma_1, \dots, \gamma_{2g}\}$ es un conjunto generador de $\pi_1(S_g, p)$ (ver Figura 5.1). Además, por la correspondencia

mencionada anteriormente, se puede hablar del número de intersección geométrico (o algebraico) de dos clases de conjugación como el número de intersección geométrico (o algebraico) de dos clases isotopía libre.

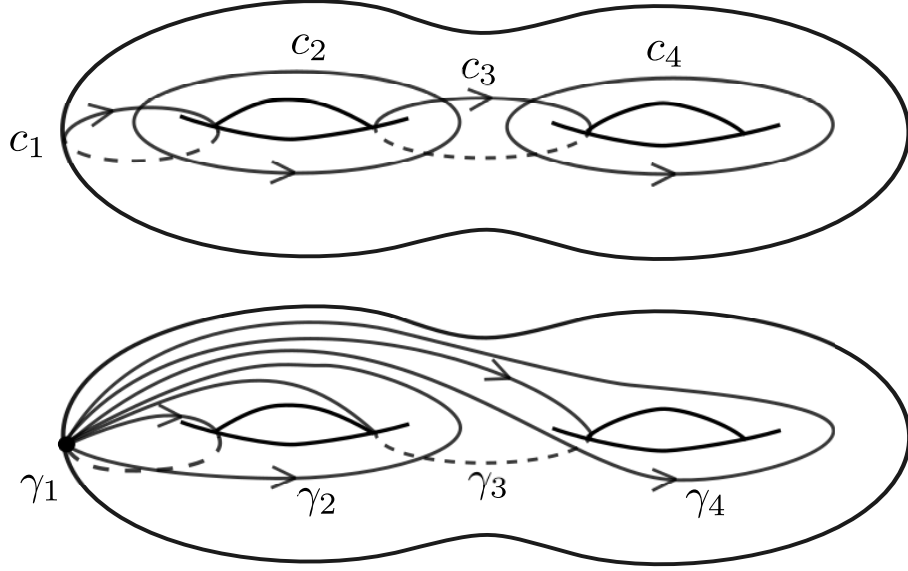


Figura 5.1: En la figura de arriba hay una cadena en una superficie de género 2 y en la figura de abajo se ven sus representantes en el grupo fundamental

Dado que Φ es un automorfismo de $\pi_1(S_g, p)$, éste actúa en el conjunto de clases de conjugación de $\pi_1(S_g, p)$. Se denotará por $\Phi(c_i)$ a la clase isotopía correspondiente a la clase de conjugación de $\Phi(\gamma_i)$.

Ahora se demostrará que $\{\Phi(c_1), \dots, \Phi(c_{2g})\}$ también es una cadena de clases de isotopía libre de curvas cerradas simples y que el número de intersección algebraico $\hat{i}(\Phi(c_i), \Phi(c_{i+1}))$ es para todos $+1$ o para todos -1 . Para esto se probarán los siguientes puntos.

1. $\Phi(c_i)$ tiene un representante simple para cada $i = 1, \dots, 2g$.
2. $i(\Phi(c_i), \Phi(c_j)) = 0$ para $|i - j| > 1$.
3. $i(\Phi(c_i), \Phi(c_{i+1})) = 1$ para cada $i = 1, \dots, 2g - 1$.

4. $\hat{i}(\Phi(c_i), \Phi(c_{i+1}))$ no depende de i .

Para probar el punto 1, obsérvese que la clase de conjugación de un elemento primitivo del $\pi_1(S_g, p)$ tiene un representante simple si y sólo si cada par de representantes de la clase de conjugación no están ligados en el infinito (ver la demostración de la Proposición 1.39). Nótese que γ_i es un elemento primitivo y que su clase de conjugación c_i tiene un representante simple, entonces cuales quiera dos representantes de la clase de conjugación de γ_i *no están* ligados en el infinito. El Lema 4.5 garantiza que ocurre lo mismo con la clase de conjugación de $\Phi(\gamma_i)$, es decir, cualquier par de representantes de esta clase de conjugación no están ligados en el infinito y por tanto $\Phi(c_i)$ tiene un representante simple.

Similarmente al paso anterior, para el paso 2 se usará el hecho de que el número de intersección geométrica de dos clases de conjugación es 0 si y sólo si cualquier par de representantes de las clases de conjugación no están ligados en el infinito (hay que demostrar este hecho). El Lema 4.5 asegura que Φ preserva esta propiedad, y por tanto se satisface el inciso 2.

La siguiente declaración es una caracterización de cuando el número de intersección geométrica de dos clases de conjugación es 1.

La intersección geométrica de dos clases de conjugación a y b es 1 si y sólo si para algún representante α de a que está ligado en el infinito con un representante dado β de b , el conjunto de representantes de a que están ligados en el infinito con β es el conjunto $\{\beta^k \alpha \beta^{-k} \mid k \in \mathbb{Z}\}$.

Una vez más, esta propiedad es invariante bajo Φ , y por tanto se cumple 3.

Para el paso 4, es suficiente notar que dadas tres clases de conjugación a , b y c tales que $i(a, b) = i(b, c) = 1$ y $i(a, c) = 0$, la coincidencia del signo de $\hat{i}(a, b)$ con el de $\hat{i}(b, c)$ se puede caracterizar en términos de las propiedades que se preservan por

Φ . Sean α, β y γ representantes arbitrarios de a, b y c respectivamente, entonces los ejes de α y γ son disjuntos y el eje de β intersecta al eje de α y al eje de γ una sola vez. Ahora, nótese lo siguiente:

Con las características de arriba, $\hat{i}(a, b)$ tiene el mismo signo que $\hat{i}(b, c)$ si y sólo si los ejes de $\alpha\beta\alpha^{-1}$ y $\gamma\beta\gamma^{-1}$ se encuentran en lados diferentes del eje de β .

Reemplazando a, b y c con c_i, c_{i+1} y c_{i+2} , aplicando el Lema 4.5 y el Corolario 4.7 a esta propiedad se demuestra 4. Así, se ha demostrado que $\{\Phi(c_i)\}$ es una cadena de clases de isotopías de curvas cerradas simples.

Por el principio de cambio de coordenadas (Teorema 1.53), existe un homeomorfismo $\phi : S_g \rightarrow S_g$ que deja fijo al punto base de $\pi_1(S_g, p)$ y satisface que $\phi(c_i) = \Phi(c_i)$ (con orientación) para cada i . Aquí c_i y $\Phi(c_i)$ denotan curvas cerradas simples representantes de sus clases de isotopía libre. Esto se puede extender a las clases de isotopía libre de curvas cerradas simples y orientadas de S_g , de tal forma que $\phi_*(c_i) = \Phi(c_i)$ para cada $i = 1, 2, \dots, 2g$. Nótese que en esto último, ϕ_* y Φ denotan las acciones inducidas en (las clases de conjugación de) los elementos de $\pi_1(S_g, p)$.

Para completar la demostración del Teorema 5.10 probaremos que la clase de isotopía $[\phi]$, actuando en $\pi_1(S_g, p)$, induce el automorfismo externo $[\Phi]$, es decir, $\phi_*^{-1} \circ \Phi$ es un automorfismo interno. En otras palabras, se debe probar que existe un $\alpha \in \pi_1(S_g, p)$ tal que para todo $\beta \in \pi_1(S_g, p)$,

$$\phi_*^{-1} \circ \Phi(\beta) = I_{\alpha^{-1}}(\beta) = \alpha^{-1}\beta\alpha.$$

Puesto que $\pi_1(S_g, p)$ es generado por $\{\gamma_1, \dots, \gamma_{2g}\}$, es suficiente probar que existe un automorfismo interno I_α de $\pi_1(S_g, p)$ tal que

$$I_\alpha \circ \phi_*^{-1} \circ \Phi(\gamma_i) = \gamma_i$$

para cada $i = 1, 2, \dots, 2g$.

Nótese que en general no es cierto que si un automorfismo de un grupo fija las clases de conjugación de cada generador, entonces es un automorfismo interno. Como un ejemplo, tómese el grupo libre generado por $\{x, y, z\}$ y considérese el automorfismo dado por $x \mapsto yxy^{-1}$, $y \mapsto y$, y $z \mapsto z$.

Se usará el hecho de que particularmente los representantes γ_i de c_i forman una cadena en el sentido de que los levantamientos de γ_i y γ_{i+1} están ligados en el infinito para cada i . Esto es consecuencia de que γ_i y γ_{i+1} se intersectan en la superficie; con más precisión, si se toma una vecindad cerrada muy pequeña en el punto p , γ_i y γ_{i+1} están ligados en la frontera de esta vecindad. Levantamientos arbitrarios de c_i y c_{i+1} podrían estar o no ligados en el infinito.

Denótese a $\phi_*^{-1} \circ \Phi$ por F . Obsérvese que F aún preserva las ligaduras en el infinito pues ϕ_* es un automorfismo de $\pi_1(S_g, p)$.

Puesto que la clase (de isotopía libre) $F(c_i)$ es igual a c_i (con orientación) para todo i , en particular se tiene que $F(c_1) = c_1$, y por tanto $F(\gamma_1) = \alpha_1^{-1} \gamma_1 \alpha_1$ para algún $\alpha_1 \in \pi_1(S_g, p)$. Se sigue que

$$I_{\alpha_1} \circ F(\gamma_1) = \gamma_1.$$

Por otro lado se sabe que $F(c_2) = c_2$, lo cual implica que $F(\gamma_2) = \alpha_2 \gamma_2 \alpha_2^{-1}$ para algún $\alpha_2 \in \pi_1(S_g, p)$. De donde se sigue que $I_{\alpha_1} \circ F(\gamma_2) = \alpha_1 \alpha_2 \gamma_2 \alpha_2^{-1} \alpha_1^{-1}$ es un elemento de c_2 , es decir, un representante de la clase de conjugación de γ_2 .

De la caracterización de dos clases de conjugación con número de intersección geométrica 1, se deduce que el conjunto de representantes de la clase de conjugación de γ_2 que están ligados en el infinito con γ_1 es $\{\gamma_1^k \gamma_2 \gamma_1^{-k} : k \in \mathbb{Z}\}$.

Recuérdese que γ_2 está ligado con γ_1 y que $I_{\alpha_1} \circ F$ preserva ligaduras en el infinito, por tanto $I_{\alpha_1} \circ F(\gamma_2)$ está ligado con $\gamma_1 = I_{\alpha_1} \circ F(\gamma_1)$. De esto y de los últimos dos párrafos anteriores se deduce que $I_{\alpha_1} \circ F(\gamma_2)$ es un elemento de $\{\gamma_1^k \gamma_2 \gamma_1^{-k} : k \in \mathbb{Z}\}$,

en otras palabras $I_{\alpha_1} \circ F(\gamma_2) = \gamma_1^{-k} \gamma_2 \gamma_1^k$ para algún $k \in \mathbb{Z}$. Defínase $\alpha := \gamma_1^k \alpha_1$. Se sigue que

$$I_\alpha \circ F(\gamma_1) = I_{\gamma_1^k} \circ I_{\alpha_1} \circ F(\gamma_1) = \gamma_1$$

y

$$I_\alpha \circ F(\gamma_2) = I_{\gamma_1^k} \circ I_{\alpha_1} \circ F(\gamma_2) = \gamma_2.$$

Hasta aquí se ha probado que γ_1 y γ_2 son puntos fijos de $I_\alpha \circ F$, ahora inductivamente se probará que $I_\alpha \circ F(\gamma_i) = \gamma_i$ para cada $i \geq 3$, y por tanto I_α es el automorfismo interno deseado. En efecto, dado que γ_1 y γ_2 son puntos fijos de $I_\alpha \circ F$, se sigue que cualquier elemento de $\{\gamma_2^l \gamma_1 \gamma_2^{-l}\}$ es fijo. Dado que γ_3 está ligado con γ_2 , γ_3 se caracteriza dentro de su clase de conjugación por estar ligado con γ_2 y porque su eje en $\mathbb{H}$ se encuentra entre los ejes de $\gamma_2^l \gamma_1 \gamma_2^{-l}$ y $\gamma_2^{l+1} \gamma_1 \gamma_2^{-(l+1)}$ para un entero l particular. Por el Corolario 4.7, el eje de $I_{\gamma_1^k \alpha_1} \circ F(\gamma_3)$ se encuentra entre los ejes de $I_{\gamma_1^k \alpha_1} \circ F(\gamma_2^l \gamma_1 \gamma_2^l) = \gamma_2^l \gamma_1 \gamma_2^l$ y $I_{\gamma_1^k \alpha_1} \circ F(\gamma_2^{l+1} \gamma_1 \gamma_2^{-(l+1)}) = \gamma_2^{l+1} \gamma_1 \gamma_2^{-(l+1)}$. Por la caracterización de γ_3 , el único eje que se encuentra entre los ejes de $\gamma_2^l \gamma_1 \gamma_2^l$ y $\gamma_2^{l+1} \gamma_1 \gamma_2^{l+1}$ es el eje de γ_3 , de donde se deduce que $I_{\gamma_1^k \alpha_1} \circ F$ fija a γ_3 . Inductivamente se demuestra que $I_{\gamma_1^k \alpha_1} \circ F$ deja fijo a γ_i , para $i = 4, \dots, 2g$. En el paso inductivo para γ_i se usa que γ_{i-1} y γ_{i-2} son ambos fijos y una caracterización de γ_i dentro de su clase de conjugación similar a la de γ_3 . Así, I_α es el automorfismo interno requerido, y con esto queda demostrado el teorema de Dehn-Nielsen-Baer. $\square$

Capítulo 6

Conclusión

En este trabajo se dio una demostración del teorema de Dehn-Nielsen-Ber: Dada una superficie orientable S_g de género $g \geq 2$. El homomorfismo de grupos

$$\begin{aligned}\sigma : \text{Mod}^\pm(S_g) &\longrightarrow \text{Out}(\pi_1(S_g, p)), \\ [\phi] &\longmapsto [\phi_*],\end{aligned}$$

es un isomorfismo.

El método de Alexander, el Lema de Švarc Milnor, y el Principio de Cambio de Coordenadas, han sido piezas claves para la demostración de este teorema. El método de Alexander fue útil para probar la inyectividad de σ . El Lema de Švarc Milnor fue esencial para la demostración del Lema 4.5, en donde se utilizó, junto con el Corolario 3.9, para garantizar que cualquier automorfismo del grupo fundamental, con la métrica hiperbólica, es una cuasi-isometría. El Lema 4.5, de que las ligaduras se preservan por automorfismos, fue el ingrediente principal para probar que los automorfismos preservan las clases de isotopías de cadenas, y poder utilizar el Principio de Cambio de Coordenadas para demostrar la sobreyectividad de σ , es decir, que para todo automorfismo externo existe un homeomorfismo que lo induce.

Como se ha visto, para demostrar el Teorema de Dehn-Nielsen-Baer se requiere desarrollar una porción considerable de teoría. Se han necesitado herramientas de topología, álgebra y geometría, lo cual dota de una belleza singular a este resultado. Pero este teorema no sólo es importante porque en ella convergen estas tres disciplinas, sino porque también abre las puertas al estudio de temas más complicados, como el caso de superficies con una cantidad finita de ponchaduras. En este caso la función σ sigue siendo inyectiva pero deja de ser sobreyectiva. Aunque no se tiene el mismo resultado, esto ha servido para unir el estudio de grupos modulares de Teichmüller con grupos de automorfismos externos de grupos libres (esto se debe a que el grupo fundamental de superficies no compactas son grupos libres). La unión de estas áreas es tal que en muchas ocasiones las técnicas e ideas usadas para obtener resultados en un área se adaptan y generalizan para obtener resultados análogos en la otra.

Si el grupo fundamental no es finitamente generado, se sabe que σ es inyectiva, pero es un área de investigación activa poder describir su imagen.

Apéndice A

Geometría Hiperbólica

A.1. Plano Hiperbólico

A.1.1. Modelos de Plano Hiperbólico

Existen distintos modelos del plano hiperbólico, en esta sección se estudian dos modelos equivalentes: el modelo del semiplano superior, y el modelo del disco de Poincaré. Modelos equivalentes se pueden consultar, por ejemplo, en [13]

El modelo del **semiplano superior** es el conjunto

$$\mathbb{H} := \{x + iy \in \mathbb{C} \mid y > 0\}$$

con la métrica Riemanniana $ds_{\mathbb{H}}^2 = \frac{dx^2 + dy^2}{y^2}$. El modelo del **disco Poincaré** es el conjunto

$$\mathbb{D} = \{x + iy \in \mathbb{C} \mid |z| < 1\}$$

con la métrica Riemanniana $ds_{\mathbb{D}}^2 = \frac{2(dx^2 + dy^2)}{(1 - (x^2 + y^2))^2}$.

Definición A.1 (Distancia hiperbólica). Sea $\gamma : [a, b] \rightarrow \mathbb{H}$ un camino diferenciable a trozos. La **longitud hiperbólica** de γ se define como $l_{\mathbb{H}}(\gamma) = \int_{\gamma} ds_{\mathbb{H}}$. Dados dos

puntos $z, w \in \mathbb{H}$, la **distancia hiperbólica** $d_{\mathbb{H}}(z, w)$ entre z y w se define como

$$d_{\mathbb{H}}(z, w) := \inf\{l_{\mathbb{H}}(\gamma) \mid \gamma \text{ es un camino diferenciable a trozos de } z \text{ a } w\}.$$

De manera análoga se define la distancia hiperbólica $d_{\mathbb{D}}$ en $\mathbb{D}$. No es difícil verificar que $(\mathbb{H}, d_{\mathbb{H}})$ y $(\mathbb{D}, d_{\mathbb{D}})$ son espacios métricos.

Definición A.2 (Geodésica). Un camino diferenciable a trozos en $\mathbb{H}$ o en $\mathbb{D}$ es una *geodésica* si la longitud de cualesquiera de sus segmentos es la mínima distancia entre sus puntos finales. La siguiente proposición da una descripción de las geodésicas en $\mathbb{H}$ y en $\mathbb{D}$.

Proposición A.3.

1. Las geodésicas en $\mathbb{H}$ son líneas euclidianas verticales o bien, semicírculos euclidianos perpendiculares a $\mathbb{R}$.
2. Las geodésicas en $\mathbb{D}$ son diámetros de $\mathbb{D}$ o arcos de círculos euclidianos perpendiculares a $\mathbb{S}^1$.

La función $C : \mathbb{H} \rightarrow \mathbb{D}$ dada por $C(z) = \frac{z-i}{z+i}$ se le conoce como *transformación de Cayley*, es una función biyectiva y continua. Realizando algunos cálculos se prueba que si $\gamma : [a, b] \rightarrow \mathbb{H}$ es un camino diferenciable a trozos, entonces $l_{\mathbb{D}}(C(\gamma)) = l_{\mathbb{H}}(\gamma)$. De esto se deduce que C preserva distancias, es decir, la transformación de Cayley es una isometría entre $\mathbb{H}$ y $\mathbb{D}$.

Al conjunto $\bar{\mathbb{R}} = \mathbb{R} \cup \{\infty\}$ se le llama la *frontera en el infinito* de $\mathbb{H}$. Similarmente, $\mathbb{S}^1$ es llamado la *frontera en el infinito* de $\mathbb{D}$; obsérvese que $\mathbb{S}^1$ es la imagen de $\bar{\mathbb{R}}$ bajo la transformación de Cayley.

A.1.2. Transformaciones de Möbius

La *esfera de Riemann* es el conjunto $\bar{\mathbb{C}} := \mathbb{C} \cup \{\infty\}$ y es homeomorfa a la esfera $\mathbb{S}^2$ [5]. Una *transformación de Möbius* es una función $T : \bar{\mathbb{C}} \rightarrow \bar{\mathbb{C}}$ de la forma

$$T(z) = \frac{az + b}{cz + d},$$

donde $a, b, c, d \in \mathbb{C}$ y $ad - bc \neq 0$. En esta parte se definen $T(\infty) := \frac{a}{b}$ y $T(\frac{-c}{d}) := \infty$.

Las transformaciones de Möbius son biyectivas; el inverso de T es $T^{-1}(z) = \frac{dz - b}{-cz + a}$ el cual también es una transformación de Möbius. Más aún, la composición de dos transformaciones de Möbius también es una transformación de Möbius. Por tanto el conjunto de transformaciones de Möbius, denotado por $\text{Möb}(\bar{\mathbb{C}})$, es un grupo con la composición usual de funciones.

Grupos lineales

El **grupo general lineal**, denotado por $GL(2, \mathbb{C})$, es el grupo de las matrices de 2×2 con coeficientes en los complejos y determinante distinto de cero, con la multiplicación usual de matrices.

El **grupo especial lineal**, denotado por $SL(2, \mathbb{C})$ es el subgrupo de $GL(2, \mathbb{C})$ que consiste de las matrices en $GL(2, \mathbb{C})$ con determinante igual a 1. Nótese que si $A \in GL(2, \mathbb{C})$ entonces $(\det A)^{-1/2}A$ es un elemento de $SL(2, \mathbb{C})$.

El **grupo proyectivo especial lineal**, denotado por $PSL(2, \mathbb{C})$, es el cociente de $SL(2, \mathbb{C})$ con el subgrupo $\{\pm I\}$.

El **grupo especial unitario**, denotado por $SU(2)$, es el grupo que consiste las matrices $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ en $SL(2, \mathbb{C})$ tales que $d = \bar{a}$ y $c = -\bar{b}$.

Se puede definir un homomorfismo de grupos $GL(2, \mathbb{C}) \rightarrow \text{Möb}(\bar{\mathbb{C}})$, tal que a cada matriz $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ se le asocia la transformación de Möbius $T(z) := \frac{az+b}{cz+d}$. Multiplicando cualquier matriz por un número complejo distinto de cero, la transformación

de Möbius correspondiente es la misma. De esto se deduce fácilmente que existe un homomorfismo sobreyectivo $SL(2, \mathbb{C}) \rightarrow \text{Möb}(\bar{\mathbb{C}})$ cuyo kernel es $\{\pm I\}$. Por el primer teorema de isomorfismos se deduce que $PSL(2, \mathbb{C}) \cong \text{Möb}(\bar{\mathbb{C}})$.

Se denota por $\text{Möb}(\mathbb{H})$ al grupo que consiste de todas las transformaciones de Möbius que preservan a $\mathbb{H}$, es decir, $\text{Möb}(\mathbb{H}) := \{T \in \text{Möb}(\bar{\mathbb{C}}) \mid T(\mathbb{H}) = \mathbb{H}\}$. Con unas cuentas fáciles se puede probar que $\text{Möb}(\mathbb{H}) \cong \text{PSL}(2, \mathbb{R})$.

De forma similar, se puede probar que $\text{Möb}(\mathbb{D}) \cong \text{PSU}(2) := \text{SU}(2)/\{\pm I\}$, el subgrupo de $\text{Möb}(\bar{\mathbb{C}})$ que deja invariante a $\mathbb{D}$.

Un cálculo simple muestra que si $T \in \text{Möb}(\mathbb{H})$ y $\gamma : [a, b] \rightarrow \mathbb{H}$ es un camino diferenciable a trozos, entonces $l_{\mathbb{H}}(T(\gamma)) = l_{\mathbb{H}}(\gamma)$. Por tanto los elementos de $\text{Möb}(\mathbb{H})$ son isometrías sobre $\mathbb{H}$; la declaración análoga para $\mathbb{D}$ también se cumple. Más aún, si se denota por $\text{Isom}^+(X)$ al grupo de isometrías de un espacio métrico X , se tiene lo siguiente:

Proposición A.4. $\text{Möb}(\mathbb{H}) = \text{Isom}^+(\mathbb{H})$ y $\text{Möb}(\mathbb{D}) = \text{Isom}^+(\mathbb{D})$.

A.1.3. Clasificación de Isometrías de $\mathbb{H}$

Definición A.5. Dos elementos T y $\hat{T}$ en $\text{Möb}(\bar{\mathbb{C}})$ son **conjugados** si existe un elemento $S \in \text{Möb}(\bar{\mathbb{C}})$ tal que $\hat{T} = STS^{-1}$.

Proposición A.6. Sean $T, \hat{T}$ y S como en la definición anterior. Si z_0 es un punto fijo de T , es decir $T(z_0) = z_0$, entonces $S(z_0)$ es un punto fijo de $\hat{T}$.

La demostración es sólo evaluar $S(z_0)$ en $\hat{T}$ y verificar que es un punto fijo.

Definición A.7. Si $T = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in GL(2, \mathbb{C})$, entonces la **traza** de T se define como

$$\text{tr}(T) = a + d.$$

Si $T \in \text{Möb}(\bar{\mathbb{C}}) \cong \text{PSL}(2, \mathbb{C})$, entonces la traza de T esta definida hasta una multiplicación por ± 1 . Por tanto, en $\text{Möb}(\bar{\mathbb{C}})$ únicamente está bien definido $\text{tr}^2(T) = (a+d)^2$. Así se tiene una función $\text{tr}^2 : \text{PSL}(2, \mathbb{C}) \rightarrow \mathbb{C}$ el cual es invariante bajo conjugación.

Ahora considérese $\text{Möb}(\mathbb{H}) = \text{PSL}(2, \mathbb{R})$. Sea $T = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$, $\det T = 1$, $\text{tr}(T) \in \mathbb{R}$. Los puntos fijos de $T \in \text{Möb}(\mathbb{H})$ son las soluciones de la ecuación $z = \frac{az+b}{cz+d}$ y se pueden calcular mediante la fórmula

$$z = \frac{(a-d) \pm \sqrt{\text{tr}^2(T) - 4}}{2c}.$$

Puesto que los coeficientes a, b, c, d son valores reales, se puede clasificar a T de las siguientes formas:

- (i) Si $\text{tr}(T) = \pm 2$ y T no es la identidad, se dice que T es una transformación **parabólica**. En este caso T tiene un único punto fijo y está en la frontera de $\mathbb{H}$.
- (ii) Si $|\text{tr}(T)| > 2$, se dice que T es una transformación **hiperbólica**. En este caso, T tiene dos puntos fijos en la frontera de $\mathbb{H}$, y por tanto fija a una única geodésica de $\mathbb{H}$.
- (iii) Si $|\text{tr}(T)| < 2$, se dice que T es una transformación **elíptica**. En este caso T tiene dos puntos fijos que son conjugados, por tanto, exactamente uno de estos puntos está en $\mathbb{H}$.

Si $T \in \text{Möb}(\mathbb{D})$, conjugando por la transformación de Cayley se puede ver que los tres casos son análogos. La única diferencia es que las transformaciones parabólicas e hiperbólicas tienen sus puntos fijos en la frontera el infinito de $\mathbb{D}$ mientras que las transformaciones elípticas tienen exactamente un punto fijo en el interior de $\mathbb{D}$, y el otro fuera.

Apéndice B

Grupo Fundamental y Espacios Cubrientes

El objetivo de este apéndice es mostrar que el grupo fundamental de una superficie se encuentra en biyección con el grupo de transformaciones de cubierta de su cubriente universal. Para esto se recopilan resultados básicos de la teoría del grupo fundamental y de la teoría de cubrientes. Los detalles de los resultados que se enuncian en esta sección, se pueden encontrar, por ejemplo, en [8] y en [11].

B.1. Grupo Fundamental

Para cualquier espacio X y cualquier punto $x \in X$, se definirá un grupo, llamado el *grupo fundamental* de X basado en x , y denotado por $\pi_1(X, x)$. Usualmente la elección del punto x carece de importancia, por esto, en muchas ocasiones es omitido en la notación.

B.1.1. Caminos y Espacios Arco-conexos

Definición B.1. Un **camino** o **arco** en X es una función continua $f : I \rightarrow X$. El punto $f(0)$ se llama **origen** del camino y el punto $f(1)$ el **final**. Se dice que f une a $f(0)$ y a $f(1)$, y que f es un camino de $f(0)$ a $f(1)$.

Obsérvese que el camino es la función f y no la imagen $f(I)$; para enfatizar esta diferencia, en este apéndice a la imagen $f(I)$ se le llamará una **curva** de X .

Ejemplo B.2. El ejemplo más sencillo de camino es el camino constante $\mathcal{C}_x : I \rightarrow X$, definido por $\mathcal{C}_x(t) = x$ para todo $t \in I$, donde x es un punto de X .

Lema B.3. Sean f y g dos caminos en X , entonces

- a) $\bar{f} : I \rightarrow X$ dado por $\bar{f}(t) = f(1 - t)$ para todo $t \in I$, es un camino de X .
- b) Si el punto final de f coincide con el origen de g , la función $f * g : I \rightarrow X$ definida por

$$f * g(t) = \begin{cases} f(2t) & \text{si } t \in [0, 1/2], \\ g(2t - 1) & \text{si } t \in [1/2, 1] \end{cases}$$

es un camino de X .

Definición B.4. Se dice que un espacio X es **arco-conexo** si dados dos puntos cualesquiera x_0, x_1 de X existe un camino en X de x_0 a x_1 .

Ejemplo B.5. $\mathbb{R}^n$ con la topología Euclidiana es arco conexo. La razón es que dados dos puntos $a, b \in \mathbb{R}^n$, la función $f : I \rightarrow \mathbb{R}^n$ definida por $f(t) = tb + (1 - t)a$ es un camino de a a b .

Ejemplo B.6. Toda n -variedad conexa es arco-conexa.

B.1.2. Homotopía de Funciones Continuas

Definición B.7. Sean X y Y espacios topológicos. Se dice que dos funciones continuas $f_0, f_1 : X \rightarrow Y$ son **homotópicas** si existe una función continua $F : X \times I \rightarrow Y$

tal que para toda $x \in X$ se cumple que $F(x, 0) = f_0(x)$ y $F(x, 1) = f_1(x)$.

La aplicación F se llama **homotopía** entre f_0 y f_1 . Se denota por $f_0 \simeq f_1$ o $F : f_0 \simeq f_1$.

Definición B.8. Un espacio X se dice que es **contráctil** si la función identidad de X es homotópica a una función constante.

La idea detrás de esta definición es que un espacio contráctil se puede deformar continuamente a un punto, como lo muestra el siguiente ejemplo:

Ejemplo B.9. El plano $\mathbb{R}^2$ es contráctil: Considérese la homotopía $F : \mathbb{R}^2 \times I \rightarrow \mathbb{R}^2$ dada por $F(x, t) = (1 - t)x$; no es difícil ver que F es una homotopía entre la función identidad de $\mathbb{R}^2$ y la función constante igual a 0. Más aún, si este ejemplo se sustituye a $\mathbb{R}^2$ por $\mathbb{R}^n$, se prueba que $\mathbb{R}^n$ es contráctil.

Por otro lado, si $f : I \rightarrow Y$ es un camino, f es homotópico al camino constante $f(0)$ mediante la homotopía $F : I \times I \rightarrow Y$ dada por $F(x, t) = f((1 - t)x)$. Cualquier función que sea homotópica a una función constante se dice que es **nulhomotópica**.

Definición B.10. Supóngase que $A \subset X$ y que f_0, f_1 son dos funciones continuas de X a Y . Se dice que f_0 y f_1 son **homotópicas relativas a A** si existe una homotopía $F : X \times I \rightarrow Y$ entre f_0 y f_1 tal que para todo $a \in A$, $F(a, t)$ no depende de t ; en otras palabras, $F(a, t) = f_0(a)$ para todo $a \in A$ y todo $t \in I$.

En este caso, a la homotopía F se le llama una **homotopía relativa a A** y se denota por $f_0 \simeq f_1(\text{rel } A)$ ó $f_0 \simeq_{\text{rel } A} f_1$. Evidentemente, para que esto pase, f_0 y f_1 deben coincidir en A . Desde luego, si $A = \emptyset$, una *homotopía relativa a A* no es otra cosa que una *homotopía*.

El siguiente lema dice que la homotopía relativa a A es una relación de equivalencia.

Lema B.11. La relación $\simeq_{rel A}$ definida en el conjunto de funciones continuas de X a Y es una relación de equivalencia.

Definición B.12. Sea $f : X \rightarrow Y$ una función continua. La **clase de homotopía libre** de f es el conjunto

$$\{g : X \rightarrow Y \mid g \simeq f\}.$$

B.1.3. Producto de Caminos

En esta subsección se introduce el concepto de producto de caminos, se considera al producto de caminos hasta homotopía relativa a $\{0, 1\}$, y se menciona con qué condiciones este producto satisface los axiomas de grupo.

Definición B.13. Sean f y g dos caminos en X tales que $f(1) = g(0)$. El **producto** (o **concatenación**) de f con g es el camino $f * g$ definido en el Lema B.3.

La concatenación de caminos de X por si sola no es muy útil, pero aunada a la relación de homotopía se obtienen resultados favorables. Por esto es oportuno definir la siguiente relación en el conjunto de caminos de X .

Definición B.14. Sean f y g dos elementos del conjunto de caminos de X . Se dice que f y g son **equivalentes** si y sólo si f y g son homotópicas relativas a $\{0, 1\}$. Esta relación se denotará por $\sim$, es decir, $f \sim g$ si y sólo si f y g son homotópicas relativas a $\{0, 1\}$.

En virtud del Lema B.11 la relación $\sim$ en el conjunto de caminos de X es una relación de equivalencia. A la clase de equivalencia de un camino f bajo la relación $\sim$, se denota por $[f]$ y se le llama **clase de homotopía de caminos de f** , abreviada como *clase de homotopía de f* cuando es claro por el contexto. Nótese la diferencia entre clase de homotopía y clase de homotopía libre (Definición B.12), en el contexto de caminos, la primera se requiere que el origen y el final sean fijos, mientras que en la segunda no.

Lema B.15. *Supóngase que f_0, f_1, g_0 y g_1 son caminos de X tales que $f_0(1) = g_0(0)$ y $f_1(1) = g_1(0)$. Si $f_0 \sim f_1$ y $g_0 \sim g_1$, entonces $f_0 * g_0 \sim f_1 * g_1$.*

El lema anterior demuestra que en el conjunto de clases de homotopía existe un producto bien definido mediante

$$[f][g] = [f * g],$$

siempre que el final de f coincida con el origen de g .

Lema B.16. *Sea f, g y h tres caminos en X . Entonces se satisfacen las siguientes propiedades:*

- a) *Si $f(1) = g(0)$ y $g(1) = h(0)$, entonces $(f * g) * h \sim f * (g * h)$.*
- b) *Si f es un camino de X con origen en x y final en y , entonces $C_x * f \sim f$ y $f * C_y \sim f$. Aquí, C_x denota el camino constante definido en el Ejemplo B.2.*
- c) *Si f como en el inciso anterior, entonces $f * \bar{f} \sim C_x$ y $\bar{f} * f \sim C_y$. Aquí, $\bar{f}$ es el camino definido en el inciso a) del Lema B.3.*

B.1.4. Grupo Fundamental

En la sección anterior se vio que el conjunto de clases de equivalencia de caminos de un espacio X prácticamente satisfacen los axiomas de un grupo. El problema es que la multiplicación no siempre está definida y el elemento neutro “no es fijo”. La manera de solucionar estos problemas es usando el concepto de camino cerrado.

Definición B.17. Un camino f de X es **cerrado** si $f(0) = f(1) = x$. En este caso se dice que f es un camino con **punto base** en x . Algunos autores usan la palabra *lazo* en lugar de camino cerrado.

Obsérvese que el producto $f * g$ está bien definido para todo par de caminos cerrados con punto base en algún $x \in X$. Se denota por $\pi_1(X, x)$ el conjunto de clases de equivalencia de caminos cerrados con punto base $x \in X$. Este conjunto posee un producto definido por $[f][g] = [f * g]$ para todo $[f], [g] \in \pi_1(X, x)$, que está bien definido en virtud del Lema B.15. El resultado siguiente establece que $\pi_1(X, x)$ con la concatenación de caminos es un grupo; al cual se le llama **grupo fundamental** de X con punto base en x .

Teorema B.18. $\pi_1(X, x)$ es un grupo.

Demostración. El resultado se deduce de la Sección B.1.3. El Lema B.15 demuestra que el producto en $\pi_1(X, x)$ está bien definido; mientras que la asociatividad, la existencia del elemento neutro, y la existencia de inversos están garantizados por el Lema B.16.

□

En vista de la asociatividad del producto de clases de equivalencia de caminos, se escribe a menudo $[f * g * h]$ en lugar de $[f * (g * h)]$. Obsérvese, sin embargo, que no podemos escribir $f * g * h$ en lugar de $f * (g * h)$.

Si se eligen dos puntos base distintos $x, y \in X$, *a priori* no existe ninguna razón por la que $\pi_1(X, x)$ y $\pi_1(X, y)$ deban estar relacionados. Sin embargo, si existe un camino de x a y entonces estos dos grupos están relacionados.

Ejemplo B.19. $\pi_1(\mathbb{S}^1, 1) \cong \mathbb{Z}$. Existe un isomorfismo $\pi_1(\mathbb{S}^1, 1) \rightarrow \mathbb{Z}$ tal que a cada clase de homotopía de un camino $f : I \rightarrow \mathbb{S}^1$ le asigna el *grado* de f . El grado, burdamente, es el número de veces en el que f *enrolla* al intervalo $[0, 1]$ sobre $\mathbb{S}^1$. Para más detalles, el lector puede consultar, por ejemplo, el Capítulo 16 de [8].

Teorema B.20. Sean X, Y dos espacios topológicos arco conexos. El grupo fundamental del producto $X \times Y$ es isomorfo al producto de los grupos fundamentales de X y Y .

Del teorema anterior y del Ejemplo B.19 se deduce que el grupo fundamental del toro $T = \mathbb{S}^1 \times \mathbb{S}^1$ es isomorfo a $\mathbb{Z}^2$.

Teorema B.21. Sean $x, y \in X$. Si existe un camino en X de x a y , entonces los grupos $\pi_1(X, x)$ y $\pi_1(X, y)$ son isomorfos.

Corolario B.22. Si X es un espacio arco conexo, $\pi_1(X, x)$ y $\pi_1(X, y)$ son isomorfos para todo par de puntos $x, y \in X$.

Ahora se dará una definición para espacios que tienen grupo fundamental trivial.

Definición B.23. Un espacio topológico X es **simplemente conexo** si es arco-conexo y $\pi_1(X, x) = \{e\}$ para algún (y por tanto para cualquier) $x \in X$.

Se puede probar que todo espacio contráctil tiene grupo fundamental trivial, y por lo tanto todo espacio contráctil es simplemente conexo. El recíproco no es cierto, como en el caso de la esfera $\mathbb{S}^2$ (ver Ejemplo B.24) que es simplemente conexa pero no contráctil.

Ejemplo B.24. La esfera $\mathbb{S}^2$ es simplemente conexa. No es difícil ver que $\mathbb{S}^2$ es arco-conexa y cualquier camino cerrado se puede deformar, hasta el punto base, continuamente sobre la esfera.

Ejemplo B.25. Para todo entero positivo n , $\mathbb{R}^n$ es simplemente conexo, ya que $\mathbb{R}^n$ es contráctil.

B.2. Espacios Cubrientes

La teoría de espacios cubrientes está estrechamente relacionada con el estudio del grupo fundamental. Muchas preguntas sobre la topología de los espacios cubrientes se pueden reducir a cuestiones puramente algebraicas de los grupos fundamentales de los espacios involucrados. Esta teoría no únicamente es importante en la topología,

también está relacionada con disciplinas como la geometría diferencial, la teoría de grupos de Lie, y la teoría de superficies Riemannianas.

B.2.1. Definición y Ejemplos de Espacios Cubrientes

En esta subsección, a menos que se indique lo contrario, se supondrá que todos los espacios son *arco conexos y localmente arco conexos* (ver [8] para estas definiciones).

Definición B.26. Una **función cubriente** (o, simplemente, un **cubriente**) es una función continua $\rho : \tilde{X} \rightarrow X$ donde cada punto $x \in X$ pertenece a un conjunto abierto $V \subset X$ tal que

$$\rho^{-1}(V) = \bigcup_j U_j$$

es una unión de conjuntos abiertos U_j disjuntos dos a dos tal que, para cada j , la restricción $\rho|_{U_j} : U_j \rightarrow V$ es un homeomorfismo. Al conjunto abierto V que satisface la condición anterior se le llama **vecindad esencial**. Al par ordenado $(\tilde{X}, \rho)$ se le llama **espacio cubriente** de X y, para cada $x \in X$, el conjunto $\rho^{-1}(x)$ se le llama la **fibra** de x . A menudo, X es llamado **espacio base** y, $\tilde{X}$ **espacio total**. Cuando $\tilde{X}$ es simplemente conexo, es decir, arco-conexo con grupo fundamental trivial, se dice que $\tilde{X}$ es un **cubriente universal** de X .

Ejemplo B.27. $(\mathbb{R}, \rho)$ es un espacio cubriente de $\mathbb{S}^1$, con $\rho : \mathbb{R} \rightarrow \mathbb{S}^1$ dado por $\rho(t) = e^{2\pi i t}$ para todo $t \in \mathbb{R}$. Cualquier arco abierto de $\mathbb{S}^1$ sirve como vecindad esencial. En particular, $\mathbb{R}$ es un cubriente universal de $\mathbb{S}^1$, pues $\mathbb{R}$ es simplemente conexo.

Ejemplo B.28. Considérese las coordenadas polares (r, θ) en $\mathbb{R}^2$. Entonces el círculo unitario $\mathbb{S}^1$ está determinado por la condición $r = 1$. Para cualquier entero n , positivo o negativo, defínase la función $\rho_n : \mathbb{S}^1 \rightarrow \mathbb{S}^1$ por la ecuación

$$\rho_n(1, \theta) = (1, n\theta).$$

La función ρ_n mapea al círculo $\mathbb{S}^1$ en si mismo n veces. No es difícil ver que si $n \neq 0$, el par $(\mathbb{S}^1, \rho_n)$ es un espacio cubriente de $\mathbb{S}^1$. Una vez más, cualquier arco abierto de $\mathbb{S}^1$ es una vecindad elemental.

Ejemplo B.29. Si $(\tilde{X}, p)$ es un espacio cubriente de X y $(\tilde{Y}, q)$ es un espacio cubriente de Y , entonces $(\tilde{X} \times \tilde{Y}, p \times q)$ es un espacio cubriente de $X \times Y$, donde la función $p \times q$ está definida por $(p \times q)(x, y) = (p(x), q(y))$. Para demostrar esto, sólo es importante notar que, si U es una vecindad elemental de un punto $x \in X$ y V es una vecindad elemental de un punto $y \in Y$, entonces $U \times V$ es una vecindad elemental de $(x, y) \in X \times Y$.

Usando este último resultado y los Ejemplos B.27 y B.28 es fácil construir ejemplos de espacios cubrientes de toro $T = \mathbb{S}^1 \times \mathbb{S}^1$. En particular, el plano $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$, el cilindro $\mathbb{R} \times \mathbb{S}^1$, o el mismo toro son ejemplos de espacios cubrientes del toro.

Ejemplo B.30. Supóngase que G actúa por homeomorfismos sobre un espacio X . Si la acción es propiamente discontinua, entonces la aplicación al cociente $\rho : X \rightarrow X/G$ es una función cubriente. Por tanto (X, ρ) es un espacio cubriente de X/G .

Proposición B.31. Sea $(\tilde{X}, \rho)$ un espacio cubriente de X . Si X es un espacio arco-conexo, entonces para cualesquiera $x, y \in X$, $\rho^{-1}(x)$ y $\rho^{-1}(y)$ tienen la misma cardinalidad.

Teorema B.32. Sea $(\tilde{X}, \rho)$ un espacio cubriente de X . Se cumple que para cualquier $x \in X$, $\rho^{-1}(x)$ es un espacio discreto.

B.2.2. Levantamiento de Caminos

Definición B.33. Sea $(\tilde{X}, \rho)$ un espacio cubriente de X y $f : Y \rightarrow X$ una función continua. Un **levantamiento** de f es una función continua $\tilde{f} : Y \rightarrow \tilde{X}$ tal que

$\rho \circ \tilde{f} = f$, es decir, el diagrama

$$\begin{array}{ccc} & & \tilde{X} \\ & \nearrow \tilde{f} & \downarrow \rho \\ Y & \xrightarrow{f} & X \end{array}$$

es conmutativo.

Lema B.34 (Levantamiento de Caminos). *Sea $\rho : \tilde{X} \rightarrow X$ una aplicación cubriente. Para cada trayectoria $\omega : I \rightarrow X$ y $\tilde{x} \in \tilde{X}$ tal que $\rho(\tilde{x}) = \omega(0)$ existe un único levantamiento $\tilde{\omega}$ de ω tal que $\tilde{\omega}(0) = \tilde{x}$. Además, si ω_0 y ω_1 son dos trayectorias en X homotópicamente equivalentes, y $\tilde{x} \in \tilde{X}$ es tal que $\rho(\tilde{x}) = \omega_0(0) = \omega_1(0)$ entonces $\tilde{\omega}_0$ y $\tilde{\omega}_1$ son homotópicamente equivalentes; en particular, $\tilde{\omega}_0$ y $\tilde{\omega}_1$ tienen el mismo origen y final.*

Teorema B.35. *Sea $\rho : \tilde{X} \rightarrow X$ una función cubriente y sea $x \in X$. Si $\tilde{X}$ es simplemente conexo, entonces $\pi_1(X, x)$ está en biyección con $\rho^{-1}(x)$.*

El teorema anterior es consecuencia del Lema B.34. El lector interesado puede consultar los detalles de la demostración, por ejemplo, en [12], en donde aparece como el Teorema 54.4.

B.2.3. Transformaciones De Cubierta

Definición B.36 (Transformaciones de Cubierta). *Sea $(\tilde{X}, \rho)$ un espacio cubriente de X . Una transformación de cubierta de $(\tilde{X}, \rho)$ es un homeomorfismo $T : \tilde{X} \rightarrow \tilde{X}$ tal que $\rho \circ T = \rho$, es decir, el diagrama*

$$\begin{array}{ccc} & & \tilde{X} \\ & \nearrow T & \downarrow \rho \\ \tilde{X} & \xrightarrow{\rho} & X \end{array}$$

es conmutativo.

El conjunto de transformaciones de cubierta de un espacio cubriente $(\tilde{X}, \rho)$ de X es un grupo con la composición usual de funciones. Este grupo se denota por $D(\tilde{X}, \rho)$.

Lema B.37. *Considérese dos transformaciones de cubierta $\varphi_0, \varphi_1 \in D(\tilde{X}, \rho)$. Si existe un punto $x \in \tilde{X}$ tal que $\varphi_0(x) = \varphi_1(x)$, entonces $\varphi_0 = \varphi_1$.*

Corolario B.38. *El grupo $D(\tilde{X}, \rho)$ actúa libremente sobre $\tilde{X}$; es decir, si $\varphi \in D(\tilde{X}, \rho)$ tal que φ no es la identidad, entonces φ no tiene puntos fijos.*

B.2.4. La Acción de $\pi_1(X, x)$ en el Conjunto $\rho^{-1}(x)$

La siguiente definición es consecuencia del Lema B.34.

Definición B.39. Sea $(\tilde{X}, \rho)$ un espacio cubriente de X y un punto $x \in X$. Se define una acción $\cdot : \pi_1(X, x) \times \rho^{-1}(x) \rightarrow \rho^{-1}(x)$ tal que $\alpha \cdot \tilde{x}$ es el punto final del único levantamiento de α , garantizado por el Lema B.34, que tiene punto inicial $\tilde{x}$.

El siguiente resultado establece una conexión entre el grupo de transformaciones de cubierta y la acción de $\pi_1(X, x)$ en $\rho^{-1}(x)$.

Proposición B.40. *Para cualesquiera $T \in D(\tilde{X}, \rho)$, $\tilde{x} \in \rho^{-1}(x)$, y $\alpha \in \pi_1(X, x)$, se cumple que*

$$T(\alpha \cdot \tilde{x}) = \alpha \cdot T(\tilde{x}).$$

Teorema B.41. *Sea $(\tilde{X}, \rho)$ un cubriente universal de X . Entonces $D(\tilde{X}, \rho)$ es isomorfo a $\pi_1(X, p)$.*

Un caso particular es cuando el espacio base X es una superficie orientable S_g de género $g \geq 2$. En este caso, un cubriente universal es $\mathbb{H}$ (o bien $\mathbb{D}$). De los teoremas B.38, B.41 y de la clasificación de las isometrías de $\mathbb{H}$, se deduce que para cualquier espacio cubriente $(\mathbb{H}, \rho)$ de S_g , todos los elementos $D(\mathbb{H}, \rho)$, son isometrías hiperbólicas.

Bibliografía

- [1] Arteaga J., *Una relación entre la geometría y el álgebra (programa de Erlangen)*. Tecné, Episteme y Didaxis, N.32, 143-148 (2012). Recuperado de <http://www.scielo.org.co/pdf/ted/n32/n32a10.pdf>
- [2] Dehn M., *Papers on Group Theory and Topology*. Springer-Verlag, New York (1987). Versión en inglés traducida del alemán y con introducciones.
- [3] Epstein D. B. A., *Curves on 2-manifolds and isotopies*. Acta Math., 115:83-107, (1966). Recuperado de <https://projecteuclid.org/euclid.acta/1485889458>
- [4] Farb B., Margalit D., *A primer on mapping class groups*. Princeton University Press, Princeton, New Jersey (2012).
- [5] Foster O., *Lectures on Riemann Surfaces*. Editorial Board. Springer-Verlag, New York (1981).
- [6] Hatcher A., *Algebraic Topology*. Versión electrónica (2001).
- [7] Hirsch M., *Differential Topology*. Graduate Texts in Mathematics, volumen 33. Springer-Verlag, New York (1994). Versión corregida de la original de 1967.
- [8] Kosniowski C., *Topología Algebraica*. Editorial Reverté, S. A., España (1986).
- [9] Löh C., *Geometric group theory: An introduction*. Universitext. Springer International Publishing, Cham, Switzerland (2017).

- [10] Margalit D., *The mathematics of Joan Birman*. Notices of the American Mathematical Society, 341 – 353 (2019). Recuperado de <https://www.ams.org/journals/notices/201903/rnoti-p341.pdf>
- [11] Massey S., *A basic course in algebraic topology*. Springer-Verlag, New York (1991).
- [12] Munkres J., *Topología* 2.^a Edición. Pearson Educación, S.A., Madrid (2002).
- [13] Series C., *Hyperbolic geometry MA448*. Online version (2013).
- [14] Thomassen C., *The Jordan-Shönflies theorem and the classification of surfaces*. Amer. Math. Monthly, 99(2):116-130, (1992). Recuperado de https://www.maa.org/sites/default/files/pdf/upload_library/22/Ford/Thomassen116-131.pdf.