



Universidad Politécnica
de Madrid

**Escuela Técnica Superior de
Ingenieros Informáticos**



Grado en Matemáticas e Informática

Trabajo Fin de Grado

**PARADOJAS en TEORÍA de
CONJUNTOS**

Autor: David Martín-Ortega Bustamante
Tutor: Alfonso Zamora Saiz

Madrid, Enero 2026

Este Trabajo Fin de Grado se ha depositado en la ETSI Informáticos de la Universidad Politécnica de Madrid para su defensa.

Trabajo Fin de Grado

Grado en Matemáticas e Informática

Título: PARADOJAS en TEORÍA de CONJUNTOS

Enero 2026

Autor: David Martín-Ortega Bustamante

Tutor: Alfonso Zamora Saiz

Departamento de Matemática Aplicada a las TIC (DMATIC)

Escuela Técnica Superior de Ingenieros Informáticos

Universidad Politécnica de Madrid

Resumen

El objetivo del trabajo es presentar de forma unificada algunas de las paradojas clásicas ligadas al origen y desarrollo de la teoría de conjuntos, poniendo el foco tanto en su mecanismo lógico interno como en la manera en que el marco axiomático moderno (ZF/ZFC) las evita, las “desactiva” o las reformula. El hilo conductor es identificar patrones comunes —autorreferencia, intentos de formar totalidades máximas (como “el conjunto de todos los conjuntos” o “el de todos los ordinales”) y argumentos de tipo diagonal— y mostrar cómo este tipo de fenómenos obliga a precisar el lenguaje y las reglas de construcción, delimitando qué colecciones pueden considerarse “conjuntos” dentro de una teoría formal y cuáles deben tratarse de otra manera.

Para ello, el trabajo se organiza en dos partes. En primer lugar se construye el vocabulario técnico necesario (axiomas de ZF/ZFC, relaciones y funciones, noción de cardinalidad, ordinales y números aleph), de manera que el análisis posterior sea autocontenido y coherente. En la segunda parte se estudian varias paradojas representativas: las paradojas de la comprensión irrestricta (en conexión con el sistema de Frege, en particular la paradoja de Russell), las paradojas de la cardinalidad del infinito (Galileo, biyecciones “contraintuitivas”, el método diagonal y la paradoja de Cantor), y las paradojas asociadas a la buena ordenación (Burali-Forti y Banach-Tarski). El resultado pretende servir como una introducción accesible pero rigurosa a estas paradojas, mostrando cómo se insertan en los fundamentos de la matemática contemporánea.

Abstract

The aim of this work is to present, in a unified way, some of the classical paradoxes connected with the origins and development of set theory, focusing both on their internal logical mechanism and on the way the modern axiomatic framework (ZF/ZFC) avoids, “neutralizes,” or reformulates them. The guiding idea is to identify common patterns—self-reference, attempts to form maximal totalities (such as “the set of all sets” or “the set of all ordinals”), and diagonal-style arguments—and to show how phenomena of this kind force us to sharpen the language and the rules of formation, thereby clarifying which collections can count as “sets” within a formal theory and which must be treated differently.

To this end, the thesis is organized into two parts. First, the necessary technical background is introduced (axioms of ZF/ZFC, relations and functions, the notion of cardinality, ordinals, and the aleph numbers) so that the subsequent discussion is self-contained and coherent. The second part studies several representative paradoxes: Russell’s paradox (in connection with Frege’s system and the failure of unrestricted comprehension), paradoxes related to the cardinality of the infinite (Galileo’s observations, “counterintuitive” bijections, the diagonal method, and Cantor’s paradox), and paradoxes tied to well-ordering (Burali-Forti and Banach-Tarski). The overall goal is to provide an accessible yet rigorous introduction to these paradoxes, showing how they fit into the foundations of contemporary mathematics.

Índice

1. Introducción	1
2. Teoría de conjuntos: nociones y conceptos básicos	3
2.1. Introducción a conjuntos	3
2.1.1. Propiedades de conjuntos	4
2.1.2. Axiomas de la teoría de conjuntos	5
2.1.3. Operaciones elementales de conjuntos	7
2.2. Relaciones, funciones y ordenaciones	8
2.2.1. Relaciones	8
2.2.2. Funciones	10
2.2.3. Ordenaciones	12
2.3. Números naturales	13
2.3.1. Propiedades	14
2.3.2. Recursión	15
2.4. Conjuntos finitos, numerables y no numerables	17
2.4.1. Cardinalidad de conjuntos	17
2.4.2. Conjuntos finitos, numerables y no numerables	18
2.5. Números cardinales	20
2.5.1. Aritmética cardinal	20
2.5.2. La cardinalidad del continuo	21
2.6. Números ordinales	25
2.6.1. Conjuntos bien ordenados	25
2.6.2. Números ordinales	25
2.6.3. Aritmética ordinal	28
2.7. Cardinales y ordinales: números aleph	30
2.7.1. Ordinales iniciales	31
2.7.2. Aritmética de los alephs	32
3. Paradojas matemáticas y su relación con la teoría de conjuntos	33
3.1. Paradojas de la comprensión irrestricta	34
3.1.1. El sistema lógico de Frege	34
3.1.2. La paradoja de Russell	35
3.2. Paradojas de la cardinalidad del infinito	37
3.2.1. Galileo y los cuadrados perfectos	37
3.2.2. «Lo veo, pero no lo creo»	38
3.2.3. La paradoja de Cantor	43

ÍNDICE

3.2.4. La paradoja del hotel infinito de Hilbert	44
3.2.5. La Hipótesis del Continuo	47
3.3. Paradojas de la buena ordenación	47
3.3.1. La paradoja de Burali-Forti	48
3.3.2. La paradoja de Banach-Tarski	49
4. Conclusiones	53
4.1. Conclusiones del trabajo	53
4.2. Análisis de impacto	54
Bibliografía	57
Anexos	59
Informe de originalidad (Turnitin)	59

Capítulo 1

Introducción

La **teoría de conjuntos** se ha consolidado como el lenguaje común de la matemática moderna. En ella se formulan de manera uniforme los **objetos**, las **relaciones** y las **construcciones** que aparecen en disciplinas tan diversas como el análisis, el álgebra o la topología. Esta unificación descansa, sin embargo, sobre una historia fundacional marcada por tensiones. La **teoría ingenua**, desarrollada a partir de los trabajos de Cantor, asumía que toda propiedad determina un conjunto bien definido; sin embargo, a comienzos del siglo XX, esta idea condujo a una serie de **paradojas** que pusieron en cuestión la consistencia misma de las matemáticas. La respuesta fue el paso a una **teoría axiomática**, en particular el sistema de **Zermelo–Fraenkel** con **Axioma de Elección (ZFC)**, que restringe cuidadosamente los principios de formación de conjuntos.

El **objetivo** de este trabajo es ofrecer una presentación unificada de algunas de las **paradojas clásicas** de la teoría de conjuntos y del infinito, analizando tanto su mecanismo interno como la forma en que el **marco axiomático** moderno las desactiva o reformula. No se trata únicamente de recopilar ejemplos llamativos, sino de estudiar sus **patrones lógicos** comunes: fenómenos de **autorreferencia**, intentos de formar **totalidades máximas** (como el conjunto de todos los conjuntos o de todos los ordinales) o construcciones que ponen en tensión nuestras intuiciones sobre **tamaño**, **medida** y **descomposición** de objetos infinitos. La idea directriz es doble: por un lado, mostrar cómo estos ejemplos motivan la introducción de axiomas como **Separación**, **Reemplazo**, **Regularidad** y **Elección**; por otro, ilustrar el alcance y los límites de ZFC como teoría de referencia para los fundamentos de la matemática.

El Capítulo 2 fija el **marco teórico** en el que se desarrollan los contenidos posteriores. Se introduce la teoría axiomática de Zermelo–Fraenkel, con y sin Axioma de Elección (Subsección 2.1.2), y se sistematizan las nociones básicas de **relaciones**, **funciones** y **ordenaciones** (Sección 2.2). A continuación se construyen los **números naturales** a partir del conjunto vacío, junto con los principios de **inducción** y **recursión** (Sección 2.3), que permiten desarrollar la aritmética elemental. A esto le sigue la noción de **cardinalidad** mediante **equipotencia**, y se distinguen los **conjuntos finitos**, **numerables** y **no numerables** (Sección 2.4), culminando en la **no numerabilidad de los reales** y la identificación de la **car-**

dinabilidad del continuo con $2^{\aleph_0}$ (Subsección 2.5.2). El capítulo se completa con la teoría de los **números ordinales** (Sección 2.6), la **aritmética ordinal** (Subsección 2.6.3) y la introducción de los **números aleph** como representantes canónicos de los **cardinales infinitos** (Sección 2.7), junto con las reglas básicas de su aritmética (Subsección 2.7.2). En conjunto, este capítulo proporciona el vocabulario y las herramientas que se utilizarán en el análisis de las paradojas.

El Capítulo 3 está dedicado específicamente a las **paradojas relacionadas con la teoría de conjuntos**, organizadas en tres bloques. En primer lugar, se abordan las **paradojas de la comprensión irrestricta** y su trasfondo lógico en el sistema de Frege (Subsección 3.1.1). A partir de la noción de clase y del principio de comprensión irrestricta se obtiene el **conjunto de Russell** (Subsección 3.1.2) y se muestra cómo su existencia conduce a una contradicción (Paradoja 1). Este análisis motiva la distinción entre **conjuntos** y **clases propias** y justifica el papel del **Esquema de separación** de Zermelo (Axioma 3) como restricción a la formación de conjuntos. En segundo lugar, se presentan las llamadas **paradojas de la cardinalidad del infinito** (Sección 3.2): el ejemplo de Galileo sobre los **cuadrados perfectos**, las **biyecciones** entre intervalos y productos cartesianos como $[0, 1]^2$ y $[0, 1]^n$, el **método diagonal** de Cantor y su uso tanto en la demostración de la no numerabilidad del continuo como en la formulación de la **paradoja de Cantor** (Subsección 3.2.2 y Paradoja 6). Dentro de este bloque se incluye también el **hotel infinito de Hilbert** (Subsección 3.2.4), que sirve para visualizar la aritmética de los **cardinales numerables**, y una discusión de la **Hipótesis del Continuo** y su estatuto de independencia respecto de ZFC. Por último, el tercer bloque reúne las **paradojas relacionadas con la buena ordenación** (Sección 3.3). Se expone la **paradoja de Burali-Forti** (Paradoja 9), que muestra que la **colección de todos los ordinales** no puede ser un conjunto, y se analiza la **paradoja de Banach-Tarski** en $\mathbb{R}^3$ (Paradoja 10), donde una **bola maciza** puede descomponerse en un número finito de piezas que se reensamblan mediante movimientos rígidos para obtener **dos copias congruentes** de la bola original. En ambos casos, se pone de relieve cómo ZFC reinterpreta estas colecciones como **clases propias** o conjuntos no medibles y qué papel desempeña el **Axioma de Elección** en la aparición de fenómenos geoméricamente contraintuitivos.

El trabajo se apoya, desde el punto de vista teórico, en el tratamiento estándar de la teoría de conjuntos de **Hrbacek y Jech** [11], mientras que el análisis de las **paradojas** y de su trasfondo histórico se basa en una selección de fuentes clásicas y contemporáneas de la literatura. El resultado pretende ser una introducción accesible pero precisa a las paradojas de la teoría de conjuntos, dirigida a lectores que ya poseen una formación básica en matemáticas y lógica, y que buscan entender cómo estas paradojas se insertan en el marco axiomático de la teoría de conjuntos y en los fundamentos de la matemática contemporánea.

Capítulo 2

Teoría de conjuntos: nociones y conceptos básicos

La **teoría de conjuntos** constituye el fundamento formal de la matemática contemporánea. En ella se formalizan los **objetos**, las **relaciones** y las **operaciones** con las que se construyen las demás estructuras. Su desarrollo permitió unificar resultados elementales y teorías abstractas bajo un mismo **marco axiomático**.

En este capítulo trabajaremos en el sistema de **Zermelo–Fraenkel** (ZF) y, cuando sea pertinente, en **ZFC** (ZF más el Axioma de Elección). Siguiendo, en lo esencial, el trabajo de Karel Hrbacek y Thomas Jech en *Introduction to Set Theory* [11], fijamos el vocabulario y las herramientas que se usarán en el resto del trabajo: una primera **introducción a conjuntos** (Sección 2.1); **relaciones, funciones y ordenaciones** (Sección 2.2); la construcción de los **números naturales** y los principios de **inducción y recursión** (Sección 2.3); la noción de cardinalidad (conjuntos finitos, numerables y el continuo) y la aritmética de **cardinales** (Secciones 2.4 y 2.5); los **números ordinales** y su aritmética (Sección 2.6); y los **números aleph** (Sección 2.7) como representantes canónicos de cardinales infinitos. Esta base servirá de soporte para los desarrollos posteriores del texto.

2.1. Introducción a conjuntos

Un **conjunto** es una colección de objetos bien definidos, llamados **elementos** o **miembros**. Esta idea, introducida formalmente por Georg Cantor, permite agruparlos mediante una **propiedad común** (véase [2]).

En el contexto de las matemáticas, los conjuntos que se estudian tienen como elementos entidades de naturaleza muy diversa: números, funciones, puntos o incluso otros conjuntos. En la formulación axiomática que adoptaremos, estos elementos se describen a su vez mediante conjuntos, como por ejemplo:

1. El conjunto de los divisores primos de 426
2. El conjunto de funciones continuas reales en $[0, 1]$

3. El conjunto de los números naturales menores que 20

Mediante combinaciones y operaciones entre conjuntos, es posible definir estructuras más complejas, como los números enteros, racionales o reales, y, en general, una amplia variedad de estructuras matemáticas. Sin embargo, para garantizar la coherencia interna de la teoría y evitar ambigüedades en la formación de conjuntos, es necesario precisar con rigor sus **propiedades** y las reglas que rigen sus **operaciones**.

2.1.1. Propiedades de conjuntos

Para definir conjuntos en matemáticas se emplea un lenguaje lógico-formal. Trabajaremos con variables para conjuntos (como $A, B, C, \dots$) y con fórmulas lógicas (como $\varphi, \psi, \dots$) que pueden depender de variables como x . Usaremos los siguientes símbolos básicos:

1. $x \in A$: expresa que el conjunto x es elemento del conjunto A .
2. $A = B$: expresa que los conjuntos A y B son iguales (véase la Definición 1).
3. $\neg\varphi$: negación de φ ; se lee “no φ ”.
4. $\varphi \wedge \psi$: conjunción; se lee “ φ y ψ ”.
5. $\varphi \vee \psi$: disyunción; se lee “ φ o ψ ”.
6. $\varphi \Rightarrow \psi$: implicación; se lee “si φ , entonces ψ ”.
7. $\varphi \Leftrightarrow \psi$: equivalencia lógica; se lee “ φ si y sólo si ψ ”.
8. $\forall x \varphi(x)$: cuantificador universal; se lee “para todo x , $\varphi(x)$ ”.
9. $\exists x \varphi(x)$: cuantificador existencial; se lee “existe al menos un x tal que $\varphi(x)$ ”.

Definición 1 (Extensionalidad) Dos conjuntos A y B son **iguales**, escrito como $A = B$, si y sólo si tienen exactamente los mismos elementos.

En particular, para cualesquiera conjuntos A, B, C se cumple:

1. $A = A$,
2. si $A = B$, entonces $B = A$,
3. si $A = B$ y $B = C$, entonces $A = C$.

Además, la igualdad respeta la pertenencia: si $A = B$ y $x \in A$, entonces $x \in B$. Estas propiedades se usarán de forma tácita en lo que sigue.

Definición 2 (Subconjunto) Un conjunto A es **subconjunto** de un conjunto B si todo elemento de A pertenece a B ; en tal caso, escribimos $A \subseteq B$.

A partir de esta noción, se verifica que $A \subseteq A$, que si $A = B$ entonces $A \subseteq B$ y $B \subseteq A$, y que la inclusión es transitiva: si $A \subseteq B$ e $B \subseteq C$, entonces $A \subseteq C$.

Definición 3 (Clase, conjunto y clase propia) En el uso informal de esta memoria, una **clase** es una colección de objetos especificada por una propiedad lógico-formal. Decimos que una clase C es un **conjunto** si existe un objeto A en el

universo de ZF tal que $x \in A$ si y sólo si x cumple la propiedad que define C . Si una clase no es un conjunto, la llamamos **clase propia**.

2.1.2. Axiomas de la teoría de conjuntos

La teoría axiomática de conjuntos se basa en el sistema de **Zermelo–Fraenkel** (ZF). A continuación se enuncian sus **axiomas** fundamentales, que establecen las reglas de existencia y formación de conjuntos.

Axioma 1 (Existencia) Existe un conjunto sin elementos.

Axioma 2 (Extensiónalidad) Si todo elemento de A pertenece a B y todo elemento de B pertenece a A , entonces $A = B$.

A partir de los Axiomas 1 y 2 se deduce la primera proposición de este trabajo:

Proposición 1 (Unicidad del vacío) Existe un único conjunto sin elementos. A este conjunto se le llama **conjunto vacío** y se denota por $\emptyset$.

Demostración. Por el Axioma 1 existe al menos un conjunto sin elementos A . Sea ahora B otro conjunto sin elementos. Queremos ver que necesariamente $A = B$. Como A no tiene elementos, es trivial que todo elemento de A pertenece a B . Análogamente, como B también es vacío, es trivial que todo elemento de B pertenece a A . Hemos comprobado, por tanto, que todo elemento de A pertenece a B y que todo elemento de B pertenece a A . Aplicando el Axioma 2, concluimos que $A = B$. $\square$

Axioma 3 (Esquema de separación) Sea $\varphi(x)$ una propiedad de x . Para cualquier conjunto A , existe un conjunto B tal que

$$x \in B \Leftrightarrow x \in A \text{ y } \varphi(x).$$

Este axioma permite formar subconjuntos a partir de un conjunto dado, seleccionando sólo aquellos elementos que satisfacen una propiedad determinada $\varphi(x)$.

Para comprender esto mejor, consideremos dos conjuntos A y B . Aplicando el Axioma 3 al conjunto A y a la propiedad “ $x \in B$ ” obtenemos, dentro de A , el subconjunto de los elementos que además pertenecen a B . Este conjunto está formado precisamente por los elementos que pertenecen simultáneamente a A y a B , y el Axioma del esquema de separación garantiza que se trata de un conjunto bien definido.

Definición 4 (Intersección) Sean A y B conjuntos. La **intersección** de A y B , denotada $A \cap B$, es el conjunto formado por los elementos que pertenecen simultáneamente a ambos:

$$A \cap B := \{x \in A \mid x \in B\}.$$

Así pues, el conjunto construido mediante el Axioma 3 coincide con la intersección $A \cap B$.

Capítulo 2. Teoría de conjuntos: nociones y conceptos básicos

Sea S una **familia** de conjuntos, es decir, un conjunto cuyos elementos son conjuntos (véase la Definición 6). Cuando $S \neq \emptyset$, tiene sentido hablar de su **intersección**, que denotamos por $\bigcap S$ y cuyos elementos son exactamente los comunes a todos los conjuntos de S :

$$x \in \bigcap S \text{ si y sólo si } x \in A (\forall A \in S)$$

Cuando S es no vacía, esta descripción determina (a lo sumo) un conjunto bien definido. En cambio, si $S = \emptyset$, la condición “para todo $A \in S$ ” no impone restricción alguna sobre x , y no existe en nuestra teoría un conjunto que recoja “todos los conjuntos posibles”. Por esta razón, en el marco de ZF no se define la intersección de la familia vacía, y la expresión $\bigcap \emptyset$ se deja sin significado.

El Axioma 3 no sólo garantiza la existencia de un subconjunto B de A definido por una propiedad $\varphi(x)$, sino también su unicidad: si B y B' son subconjuntos de A tales que

$$x \in B \iff x \in A \text{ y } \varphi(x), \quad x \in B' \iff x \in A \text{ y } \varphi(x),$$

entonces necesariamente $B = B'$. Así, es legítimo escribir $\{x \mid \varphi(x)\}$ (o, más precisamente, $\{x \in A \mid \varphi(x)\}$) para referirse al conjunto de elementos de A que satisfacen $\varphi(x)$; por ejemplo, $\{x \in A \mid x \in B\}$ en el caso de la intersección.

Axioma 4 (Par) Para cualesquiera conjuntos A y B , existe un conjunto C , que denotamos $\{A, B\}$, tal que $x \in C$ si y sólo si $x = A$ o $x = B$.

Axioma 5 (Unión) Para cualquier conjunto S , existe un conjunto U tal que $x \in U$ si y sólo si $x \in A$ para algún $A \in S$.

En particular, el Axioma 4 garantiza la existencia del conjunto $\{A, B\}$ que “agrupa” a A y B como elementos. Combinándolo con el Axioma 5, definimos la unión de dos conjuntos.

Definición 5 (Unión) La **unión** de A y B , denotada $A \cup B$, es el conjunto de todos los elementos x que pertenecen a A , a B , o a ambos.

Más en general, la unión de un sistema S es exactamente el conjunto de todos los x que pertenecen a algún miembro de S , esto es, $\bigcup S = \{x \mid \exists A \in S (x \in A)\}$.

Axioma 6 (Conjunto potencia) Para cualquier conjunto S , existe un conjunto P , denotado por $\mathcal{P}(S)$, tal que $A \in P$ si y sólo si $A \subseteq S$.

Los seis primeros axiomas permiten construir y combinar conjuntos a partir de otros. Sin embargo, la teoría requiere además algunos **axiomas adicionales** que amplían el marco de existencia de conjuntos, aunque su comprensión completa dependerá de conceptos que se introducirán más adelante.

Axioma 7 (Infinito) Existe un conjunto **inductivo**, es decir, un conjunto I tal que $\emptyset \in I$ y para todo $x \in I$, también $x \cup \{x\} \in I$.

En particular, a partir de este axioma se puede construir el conjunto de los **números naturales**, denotado por $\mathbb{N}$ (véase Definición 26).

Axioma 8 (Reemplazo) Sea $\varphi(x, y)$ una propiedad tal que, para todo x , existe un único y con $\varphi(x, y)$. Entonces, para cada conjunto A , existe un conjunto B tal que para todo $x \in A$, existe $y \in B$ con $\varphi(x, y)$.

Es decir, si a cada x le corresponde de manera unívoca un y , el axioma garantiza que la imagen de cualquier conjunto A mediante esta correspondencia vuelve a ser un conjunto. Por ejemplo, si tomamos $\varphi(x, y)$ como la propiedad $y = \{x\}$. Para cada x existe un único y que cumple esta propiedad, concretamente el conjunto unitario $\{x\}$. Aplicando el axioma de reemplazo, se deduce que, dado un conjunto A , existe el conjunto $B = \{\{x\} \mid x \in A\}$, formado por todos los subconjuntos unitarios de A .

Axioma 9 (Regularidad) Todo conjunto no vacío A contiene un **elemento disjunto** de A , es decir, para todo $A \neq \emptyset$ existe un $x \in A$ tal que $x \cap A = \emptyset$.

Este axioma excluye, por ejemplo, la existencia de conjuntos que se contengan a sí mismos ($x \in x$) o de cadenas infinitas descendentes de pertenencia ($\cdots \in x_2 \in x_1 \in x_0$). En términos intuitivos, afirma que todos los conjuntos están “bien fundados” y permite organizar el universo de la teoría de conjuntos en niveles contruidos de manera iterativa.

Todos los axiomas enunciados hasta ahora constituyen el núcleo de la **teoría axiomática de conjuntos de Zermelo–Fraenkel** (ZF). Este sistema proporciona un marco formal suficiente para desarrollar la mayor parte de la matemática moderna, sustituyendo la noción intuitiva de conjunto por un conjunto de reglas precisas que determinan qué conjuntos pueden existir y cómo se construyen a partir de otros.

Definición 6 (Familia de conjuntos) Una **familia de conjuntos** es un conjunto $\mathcal{F}$ cuyos elementos son, a su vez, conjuntos.

Axioma 10 (Elección) Para toda familia $\mathcal{F}$ de conjuntos no vacíos existe una función f con dominio $\mathcal{F}$ tal que $f(A) \in A$ para todo $A \in \mathcal{F}$ (véase Definición 11).

Al añadir este último axioma, obtenemos **ZFC**. En lo que sigue trabajaremos dentro de este sistema. No estudiaremos la teoría sin elección; cuando algún resultado dependa esencialmente de este axioma lo indicaremos de forma explícita.

2.1.3. Operaciones elementales de conjuntos

Introducimos las **operaciones elementales** entre conjuntos y algunas de sus **propiedades básicas**. De acuerdo con la Definición 2, escribimos $A \subseteq B$ para indicar que todo elemento de A pertenece también a B (inclusión). De aquí se siguen, de manera rutinaria, las propiedades: $A \subseteq A$; si $A \subseteq B$ y $B \subseteq A$, entonces $A = B$; y si $A \subseteq B$ y $B \subseteq C$, entonces $A \subseteq C$. Diremos que A es **subconjunto propio** de B y escribimos $A \subsetneq B$ cuando $A \subseteq B$ pero $A \neq B$; análogamente, $B \supsetneq A$ y $B \supset A$. Tras esta formalización breve, recordamos la intersección y la unión de conjuntos según las Definiciones 4 y 5. Introducimos ahora otras operaciones básicas.

Llamaremos **diferencia** de A y B , denotada $A \setminus B$, al conjunto de los elementos de A que no pertenecen a B . La **diferencia simétrica** se define por

$$A \Delta B = (A \setminus B) \cup (B \setminus A).$$

Las operaciones de unión, intersección y diferencia verifican las propiedades habituales de conmutatividad, asociatividad, distributividad y las leyes de De Morgan. En particular, para cualesquiera conjuntos A, B, C se tiene:

1. **Conmutatividad:** $A \cap B = B \cap A$, $A \cup B = B \cup A$.
2. **Asociatividad:** $(A \cap B) \cap C = A \cap (B \cap C)$, $(A \cup B) \cup C = A \cup (B \cup C)$.
3. **Distributividad:** $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$, $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$.
4. **Leyes de De Morgan:** $C \setminus (A \cap B) = (C \setminus A) \cup (C \setminus B)$, $C \setminus (A \cup B) = (C \setminus A) \cap (C \setminus B)$.

Diremos que dos conjuntos A y B son **disjuntos** si $A \cap B = \emptyset$. Más generalmente, una familia S es de conjuntos **mutuamente disjuntos** si para todo $A, B \in S$ con $A \neq B$ se cumple $A \cap B = \emptyset$.

2.2. Relaciones, funciones y ordenaciones

Las nociones de **relación**, **función** y **orden** permiten pasar de conjuntos considerados de forma aislada a estructuras en las que sus elementos interactúan. En primer lugar introduciremos el concepto de **par ordenado** (Definición 7), que permite entender el concepto de **relación binaria** como subconjunto de un producto (Definición 9). Sobre esta base, interpretaremos las funciones como relaciones que asignan a cada elemento una única imagen (Definición 11), e introduciremos las **relaciones de equivalencia** y su vínculo con **particiones** y **cocientes** (Definiciones 16 y 18). Finalmente, presentaremos las distintas nociones de **orden** (parcial y total), así como comparabilidad, cadenas y elementos extremos (Definiciones 19, 21 y 23). Estos conceptos aparecerán de forma recurrente en capítulos posteriores, en particular en la teoría de ordinales.

2.2.1. Relaciones

Para hablar de relaciones necesitamos distinguir el orden de los elementos. Si sólo consideramos el conjunto $\{a, b\}$, no hay diferencia entre escribir $\{a, b\}$ y $\{b, a\}$, ya que $\{a, b\} = \{b, a\}$. En cambio, cuando escribimos (a, b) queremos que el primer componente sea a y el segundo sea b , de modo que, en general, $(a, b) \neq (b, a)$. Para formalizar esta idea usando únicamente conjuntos, introducimos la siguiente definición.

Definición 7 (Par ordenado) Definimos el **par ordenado** (a, b) como el conjunto

$$(a, b) := \{\{a\}, \{a, b\}\}.$$

Con esta codificación, el primer componente se recupera como el único elemento de (a, b) que es un conjunto unitario, y el segundo como el elemento que aparece junto al primero. El siguiente resultado muestra que esta noción de “primer” y “segundo” componente es canónica.

2.2. Relaciones, funciones y ordenaciones

Proposición 2 (Criterio de igualdad de pares) *Dos pares ordenados son iguales si y sólo si coinciden componente a componente:*

$$(a, b) = (a', b') \iff a = a' \text{ y } b = b'.$$

Demostración. Si $(a, b) = (a', b')$, entonces $\{\{a\}, \{a, b\}\} = \{\{a'\}, \{a', b'\}\}$. En cada lado aparece exactamente un conjunto unitario, de modo que $\{a\} = \{a'\}$ y por tanto $a = a'$. Los otros elementos, $\{a, b\}$ y $\{a', b'\}$, también coinciden, y como ya sabemos que $a = a'$, se deduce $b = b'$. La implicación recíproca es inmediata. $\square$

A partir de los pares ordenados ya definidos (Definición 7), se pueden construir **tuplas ordenadas** más largas, como:

1. Triples ordenados: $(a, b, c) = ((a, b), c)$,
2. Cuádruplas ordenadas: $(a, b, c, d) = ((a, b, c), d)$,

y análogamente para aquellas de mayor longitud.

Definición 8 (Producto cartesiano) *Dados conjuntos A y B , el **producto cartesiano** de A y B es el conjunto*

$$A \times B := \{ (a, b) \mid a \in A, b \in B \},$$

formado por todos los pares ordenados cuyo primer componente pertenece a A y cuyo segundo componente pertenece a B .

Definición 9 (Relación binaria) *Sean A y B conjuntos. Una **relación binaria** entre A y B es un subconjunto $R \subseteq A \times B$. Si $R \subseteq A \times A$, decimos que R es una relación en A . Cuando $(a, b) \in R$, solemos escribir también $a R b$.*

Para una mejor comprensión de la definición anterior, veamos un ejemplo. Sea A un conjunto y consideremos su conjunto potencia $\mathcal{P}(A)$. Definimos una relación $R \subseteq \mathcal{P}(A) \times \mathcal{P}(A)$ escribiendo $R := \{ (X, Y) \in \mathcal{P}(A) \times \mathcal{P}(A) \mid X \subseteq Y \}$. Entonces $X R Y$ significa exactamente que X es un subconjunto de Y . En otras palabras, la relación “estar incluido en” se puede ver como una relación binaria sobre $\mathcal{P}(A)$.

Dada una relación binaria $R \subseteq A \times B$, su **relación inversa** de R es el subconjunto $R^{-1} \subseteq B \times A$ definido por

$$R^{-1} := \{ (b, a) \mid (a, b) \in R \}.$$

De este modo, $b R^{-1} a$ significa precisamente que $a R b$.

Definición 10 (Composición de relaciones) *Sean $R \subseteq A \times B$ y $T \subseteq B \times C$ relaciones binarias. La **composición** $T \circ R$ es la relación en $A \times C$ definida por*

$$T \circ R := \{ (a, c) \mid \exists b \in B ((a, b) \in R \text{ y } (b, c) \in T) \}.$$

Intuitivamente, $(a, c) \in T \circ R$ cuando es posible “ir de a hasta c ” primero mediante R y después mediante T .

2.2.2. Funciones

Una **función**, en matemáticas, es un procedimiento que asigna a cada elemento a de un conjunto (dominio) un único elemento b de otro conjunto (codominio), llamado el valor de la función en a . Una función es, por tanto, un tipo especial de relación binaria en la que cada elemento del dominio está relacionado con un único elemento del codominio. Usaremos la codificación de pares ordenados de la Definición 7 y la composición de la Definición 10.

Definición 11 (Función) Sean A y B conjuntos. Una relación $f \subseteq A \times B$ es una **función** de A en B si para todo $a \in A$ existe un único $b \in B$ tal que $(a, b) \in f$. En este caso escribimos $f : A \rightarrow B$.

Al conjunto A lo llamamos **dominio** de f , y al conjunto B lo llamamos **codominio** (dominio de llegada) de f . Para cada $a \in A$ denotamos por $f(a)$ al único $b \in B$ tal que $(a, b) \in f$. La **imagen** de f es el conjunto

$$\text{Im } f := \{ f(a) \mid a \in A \} \subseteq B.$$

Proposición 3 (Extensionalidad de funciones) Sean $f : A \rightarrow B$ y $g : A \rightarrow B$ funciones con el mismo dominio. Entonces $f = g$ si y sólo si $f(a) = g(a)$ para todo $a \in A$.

Demostración. Si $f = g$, entonces trivialmente $f(a) = g(a)$ para todo $a \in A$. Recíprocamente, si $f(a) = g(a)$ para todo $a \in A$, todo par de f es de la forma $(a, f(a))$ con $a \in A$, y como $f(a) = g(a)$, se tiene $(a, f(a)) = (a, g(a))$, de modo que $(a, f(a)) \in g$. Por tanto, $f \subseteq g$. El mismo argumento, a la inversa, muestra que $g \subseteq f$. En consecuencia, $f = g$. $\square$

Definición 12 (Sobreyectividad, inyectividad, biyectividad) Sea $f : A \rightarrow B$ una función.

1. f es **inyectiva** si para cualesquiera $a_1, a_2 \in A$, de $f(a_1) = f(a_2)$ se sigue $a_1 = a_2$.
2. f es **sobreyectiva** sobre B si $\text{Im } f = B$, es decir, si todo elemento de B es imagen de algún elemento de A .
3. f es **biyectiva** si es a la vez inyectiva y sobreyectiva.

La composición de funciones se define de la forma habitual, en términos de la aplicación sucesiva de una función tras otra.

Definición 13 (Composición de funciones) Sean $f : A \rightarrow B$ y $g : B \rightarrow C$ funciones. La **composición** de g con f es la función

$$g \circ f : A \rightarrow C, \quad (g \circ f)(a) := g(f(a)) \quad (a \in A).$$

Demostración. Para cada $a \in A$, existe un único $b = f(a) \in B$, y para ese b existe un único $c = g(b) \in C$. Por tanto, a cada $a \in A$ le corresponde un único $c \in C$, concretamente $(g \circ f)(a)$, y la definición anterior determina una función. $\square$

Para estudiar la invertibilidad de funciones necesitaremos la función identidad.

Definición 14 (Identidad e inversa) Sea A un conjunto. La **función identidad** en A es la función

$$\text{id}_A : A \rightarrow A, \quad \text{id}_A(a) := a.$$

Sea $f : A \rightarrow B$ una función. Decimos que f es **invertible** si existe una función $g : B \rightarrow A$ tal que

$$g \circ f = \text{id}_A \quad \text{y} \quad f \circ g = \text{id}_B.$$

En este caso, g se llama la **inversa** de f y se denota f^{-1} .

Teorema 1 (Invertibilidad) Sea $f : A \rightarrow B$ una función. Entonces f es invertible si y sólo si es biyectiva.

Demostración. ($\Rightarrow$) Supongamos que f es invertible. Existe entonces una función inversa $g : B \rightarrow A$ tal que $g \circ f = \text{id}_A$ y $f \circ g = \text{id}_B$. Sean $x_1, x_2 \in A$ tales que $f(x_1) = f(x_2)$. Aplicando g a ambos lados: $g(f(x_1)) = g(f(x_2)) \Rightarrow (g \circ f)(x_1) = (g \circ f)(x_2)$. Como $g \circ f = \text{id}_A$, tenemos que $x_1 = x_2$. Por tanto, f es inyectiva. Sea y un elemento cualquiera de B . Definimos $x := g(y)$. Al aplicar f : $f(x) = f(g(y)) = (f \circ g)(y)$. Como $f \circ g = \text{id}_B$, resulta que $f(x) = y$. Así, todo $y \in B$ tiene una preimagen, por lo que f es sobreyectiva. Concluimos que si f es invertible, entonces es biyectiva.

($\Leftarrow$) Supongamos que f es biyectiva (inyectiva y sobreyectiva). Definimos la función $g : B \rightarrow A$ de la siguiente manera: para cada $y \in B$, sea $g(y)$ el único elemento $x \in A$ tal que $f(x) = y$. Veamos que g es la inversa: Para todo $x \in A$, si llamamos $y = f(x)$, por definición de g tenemos que $g(y) = x$, luego $(g \circ f)(x) = x$. Para todo $y \in B$, si tomamos $x = g(y)$, por definición sabemos que $f(x) = y$, luego $(f \circ g)(y) = y$. Por tanto, f es invertible. $\square$

Definición 15 (Propiedades binarias) Sea R una relación binaria en un conjunto A (es decir, $R \subseteq A \times A$). Decimos que:

1. R es **reflexiva** en A si para todo $a \in A$, se tiene aRa .
2. R es **simétrica** en A si para todo $a, b \in A$, aRb implica bRa .
3. R es **transitiva** en A si para todo $a, b, c \in A$, si aRb y bRc , entonces aRc .
4. R es **antisimétrica** en A si para todo $a, b \in A$, si aRb y bRa , entonces $a = b$.
5. R es **asimétrica** en A si para todo $a, b \in A$, de aRb se sigue que no se cumple bRa , es decir, $\neg(bRa)$.

Definición 16 (Relación de equivalencia) Una relación E en A es una **relación de equivalencia** si es reflexiva, simétrica y transitiva.

Definición 17 (Clase de equivalencia y cociente) Sea E una relación de equivalencia en A . La **clase** de a es

$$[a]_E = \{x \in A \mid xEa\},$$

y el **conjunto cociente** es $A/E = \{[a]_E \mid a \in A\}$.

Definición 18 (Partición) Una **partición** de A es una familia S de subconjuntos no vacíos tal que sus elementos son disjuntos dos a dos y $\bigcup S = A$.

Capítulo 2. Teoría de conjuntos: nociones y conceptos básicos

Dada una partición S de A , consideramos la relación

$$E_S = \{(a, b) \in A \times A \mid \exists C \in S, a \in C \text{ y } b \in C\},$$

que relaciona dos elementos de A precisamente cuando pertenecen al mismo bloque de la partición S .

Proposición 4 (Correspondencia entre equivalencias y particiones) *Sea A un conjunto.*

1. Si E es una relación de equivalencia en A , entonces A/E es una partición de A y la relación asociada a esta partición coincide con E , es decir, $E_{A/E} = E$.
2. Si S es una partición de A , entonces E_S es una relación de equivalencia en A y el cociente asociado a E_S recupera la partición inicial, es decir, $A/E_S = S$.

Demostración. La demostración se obtiene combinando varias propiedades elementales de relaciones de equivalencia y particiones. Todos estos argumentos pueden encontrarse, por ejemplo, en [11, págs. 31–32]. $\square$

Cuando S es una partición de A , diremos que un subconjunto $X \subseteq A$ es un **conjunto de representantes** de S si para todo $C \in S$ existe $a \in C$ tal que $X \cap C = \{a\}$.

2.2.3. Ordenaciones

En esta sección trabajaremos con la siguiente notación: utilizaremos el símbolo $\leq$ para denotar una relación de orden parcial y el símbolo $<$ para denotar una relación de orden estricta.

Definición 19 (Orden parcial) Una relación $\leq$ en A es un **orden parcial** si es reflexiva, antisimétrica y transitiva (en el sentido de la Definición 15). En este caso, el par $(A, \leq)$ se llama un **conjunto ordenado**.

Definición 20 (Orden estricto) Una relación $<$ en A es un **orden estricto** si es asimétrica y transitiva.

Es habitual pasar de un orden no estricto $\leq$ a uno estricto $<$, y viceversa. Si $(A, \leq)$ es un conjunto ordenado, podemos definir

$$a < b \iff a \leq b \text{ y } a \neq b,$$

obteniendo así un orden estricto en A . Recíprocamente, si $<$ es una ordenación estricta en A , ponemos

$$a \leq b \iff a < b \text{ o } a = b,$$

y esta relación resulta ser un orden parcial. En lo que sigue pasaremos de una descripción a otra sin mencionarlo de forma explícita.

Definición 21 (Comparabilidad y orden total) Dados $a, b \in A$ y un orden $\leq$ en A , decimos que a y b son **comparables** si $a \leq b$ o $b \leq a$. El orden $\leq$ es **lineal** (total) si todo par de elementos de A es comparable; en ese caso, $(A, \leq)$ es un **conjunto linealmente ordenado**.

Un orden en A induce de manera natural órdenes en sus subconjuntos. Si $(A, \leq)$ es un conjunto ordenado y $B \subseteq A$, trabajaremos en B con el **orden inducido** por $\leq$, es decir, comparamos $a, b \in B$ escribiendo $a \leq b$ exactamente cuando la desigualdad se cumple en A . Con esta restricción, B vuelve a ser un conjunto ordenado.

Definición 22 (Cadena) Sea $B \subseteq A$, donde A está ordenado por $\leq$. Decimos que B es una **cadena** en A si todos sus elementos son comparables entre sí.

Definición 23 (Mínimos, máximos y versiones relativas) Sea $(A, \leq)$ un conjunto ordenado y $B \subseteq A$. Un $b \in B$ es **mínimo** si $b \leq x$ para todo $x \in B$ (y **máximo** si $x \leq b$ para todo $x \in B$). Es **mínimo relativo** si no existe $x \in B$ con $x \neq b$ y $x \leq b$ (análogamente para **máximo relativo**).

Obsérvese que un subconjunto B tiene a lo sumo un mínimo: si hubiera dos, al ser comparables entre sí, coincidirían. Todo mínimo (cuando existe) es, en particular, un mínimo relativo. Además, si B es una cadena, las dos nociones se identifican: cualquier mínimo relativo de B resulta ser también el mínimo de B . Los comentarios duales valen para máximos y máximos relativos.

Definición 24 (Ínfimo y supremo) Sea $(A, \leq)$ un conjunto ordenado y $B \subseteq A$.

1. $a \in A$ es **cota inferior** de B si $a \leq x$ para todo $x \in B$. El **ínfimo** de B es la mayor de sus cotas inferiores, cuando existe.
2. $a \in A$ es **cota superior** de B si $x \leq a$ para todo $x \in B$. El **supremo** de B es la menor de sus cotas superiores, cuando existe.

La diferencia entre el mínimo de B y una cota inferior de B es que esta última no necesita pertenecer a B . Un conjunto puede tener muchas cotas inferiores, pero el conjunto de todas ellas puede tener a lo sumo un mayor elemento; por tanto, B puede tener a lo sumo un ínfimo.

En consecuencia, si b es el mínimo de B , entonces b coincide con el ínfimo de B , y si b es el ínfimo de B y además $b \in B$, entonces b es el mínimo de B . De forma equivalente, un elemento $b \in A$ es el ínfimo de B en $(A, \leq)$ si y sólo si se verifican: (i) $b \leq x$ para todo $x \in B$, y (ii) si $b' \leq x$ para todo $x \in B$, entonces $b' \leq b$. Los enunciados duales se obtienen sustituyendo “mínimo/ínfimo” por “máximo/supremo” y “ $\leq$ ” por “ $\geq$ ”.

2.3. Números naturales

En el marco axiomático de la teoría de conjuntos, los **números naturales** se construyen a partir del conjunto vacío mediante el operador **sucesor** (Definición 25) y el **Axioma del Infinito** (Axioma 7), que garantiza la existencia de un conjunto inductivo. Identificaremos 0 con $\emptyset$ y, en general, cada natural con el conjunto de todos sus predecesores, (construcción de von Neumann, [15]). A partir de esta construcción definiremos el conjunto de los números naturales, que denotaremos por $\mathbb{N}$, y estableceremos el principio de **inducción** y el **teorema de recursión**, que fundamentan la aritmética básica y el **orden natural**.

Capítulo 2. Teoría de conjuntos: nociones y conceptos básicos

Definición 25 (Sucesor) Para cualquier conjunto x , el **sucesor** de x es

$$S(x) := x \cup \{x\}.$$

Definición 26 (Números naturales) Llamamos **conjunto de los números naturales** a un conjunto $\mathbb{N}$ que satisface:

1. $\mathbb{N}$ es inductivo.
2. Si I es un conjunto inductivo, entonces $\mathbb{N} \subseteq I$.

Es decir, $\mathbb{N}$ es el menor de los conjuntos inductivos, por lo que está contenido en todos ellos. A partir de ahora fijamos un conjunto con estas propiedades y lo denotamos por $\mathbb{N}$.

A partir de ahora, si $n \in \mathbb{N}$ es un número natural, escribiremos también

$$n + 1 := S(n)$$

para denotar su sucesor.

Proposición 5 (Existencia y minimalidad de $\mathbb{N}$) Existe un único conjunto $\mathbb{N}$ que cumple las condiciones de la Definición 26. Además, se tiene

$$\mathbb{N} = \bigcap \{ I \mid I \text{ es un conjunto inductivo} \}.$$

Demostración. Esta proposición se obtiene construyendo $\mathbb{N}$ como la intersección de todos los subconjuntos inductivos de un conjunto inductivo dado (garantizado por el Axioma 7) y comprobando que dicha intersección sigue siendo inductiva y es mínima entre todas ellas. $\square$

En la construcción de von Neumann, los primeros números naturales se describen explícitamente como

$$0 := \emptyset, \quad 1 := S(0) = 0 \cup \{0\} = \{\emptyset\}, \quad 2 := S(1) = 1 \cup \{1\} = \{\emptyset, \{\emptyset\}\},$$

y, en general,

$$n + 1 = S(n) = n \cup \{n\}.$$

Cada número natural coincide, por tanto, con el conjunto de todos sus predecesores.

Definición 27 (Orden natural) Para $m, n \in \mathbb{N}$ definimos por una parte, $m < n$, si $m \in n$, y por otra $m \leq n$, si $m < n$ o $m = n$.

2.3.1. Propiedades

Teorema 2 (Principio de inducción) Sea $P(n)$ una propiedad definida para $n \in \mathbb{N}$. Son equivalentes las siguientes afirmaciones:

1. **Inducción.** Si se verifica que $P(0)$ es verdadera, y que para todo $n \in \mathbb{N}$, si $P(n)$ es verdadera, entonces $P(n + 1)$ también lo es. Entonces $P(n)$ es verdadera para todo $n \in \mathbb{N}$.

2. **Inducción fuerte.** Si para todo $k < n \in \mathbb{N}$ se cumple que $P(k) \implies P(n)$. Entonces $P(n)$ es verdadera para todo $n \in \mathbb{N}$.

Cualquiera de estas formulaciones se denomina **principio de inducción** en $\mathbb{N}$.

Demostración. Consideremos el subconjunto $A := \{n \in \mathbb{N} \mid P(n)\}$. Las hipótesis de (1) dicen precisamente que A es inductivo: $0 \in A$ y, si $n \in A$, entonces $n + 1 \in A$. Por la Definición 26, $\mathbb{N}$ es el menor conjunto inductivo, de modo que necesariamente $\mathbb{N} \subseteq A$. En particular, $P(n)$ es verdadera para todo $n \in \mathbb{N}$.

La formulación de inducción fuerte se obtiene a partir de la anterior, definiendo la propiedad $Q(n) : (\forall k < n, P(k))$ y aplicando el principio de inducción a Q (véase [11, pág. 44] para una demostración detallada). $\square$

En el orden natural definido en la Definición 27 se tiene siempre $0 \leq n$ para todo $n \in \mathbb{N}$, y para cualesquiera $k, n \in \mathbb{N}$ se cumple

$$k < n + 1 \iff (k < n \text{ o } k = n),$$

propiedades que se demuestran directamente por inducción.

Definición 28 (Buen orden) Sea $(A, <)$ un conjunto linealmente ordenado por un orden estricto $<$. Decimos que $<$ es un **buen orden** si todo subconjunto no vacío de A tiene mínimo. En ese caso, decimos que $(A, <)$ es un **conjunto bien ordenado**.

Teorema 3 (Buen orden de $\mathbb{N}$) Con el orden natural $(<)$ definido en la Definición 27, el conjunto $(\mathbb{N}, <)$ es un conjunto bien ordenado.

Demostración. Sea $B \subseteq \mathbb{N}$ un subconjunto no vacío. Queremos ver que B tiene un mínimo. Supongamos, por absurdo, que B no tiene elemento mínimo y consideremos su complemento $C := \mathbb{N} \setminus B$. Definimos la propiedad $P(n) : n \in C$ (es decir, $n \notin B$). En primer lugar, $P(0)$ es verdadera: si $0 \notin C$, entonces $0 \in B$ y, al no haber ningún natural menor que 0, este sería el mínimo de B , en contra de la hipótesis.

Supongamos ahora que $n \in \mathbb{N}$ satisface que $P(k)$ es verdadera para todo $k < n$, es decir, que ningún $k < n$ pertenece a B . Si $n \notin C$, entonces $n \in B$ y, como ningún elemento de B es menor que n , se seguiría que n es el mínimo de B , lo cual contradice nuevamente la hipótesis inicial. Por tanto, de la validez de $P(k)$ para todo $k < n$ se deduce que también $P(n)$ es verdadera.

Hemos verificado así las hipótesis de la versión de inducción fuerte (Teorema 2), de modo que $P(n)$ es verdadera para todo $n \in \mathbb{N}$; es decir, $C = \mathbb{N}$ y, en consecuencia, $B = \emptyset$. Esto contradice que B fuera no vacío. Luego toda parte no vacía de $\mathbb{N}$ posee un mínimo, y $(\mathbb{N}, <)$ es un conjunto bien ordenado. $\square$

Como consecuencia del buen orden de $(\mathbb{N}, <)$, todo subconjunto no vacío de $\mathbb{N}$ que esté acotado superiormente posee un máximo.

2.3.2. Recursión

Para manejar definiciones iteradas, fijamos la notación de secuencias finitas e infinitas como funciones con dominio un natural o $\mathbb{N}$, respectivamente.

Capítulo 2. Teoría de conjuntos: nociones y conceptos básicos

Definición 29 (Secuencias finitas e infinitas) Una **secuencia finita** de longitud $n \in \mathbb{N}$ con valores en A es una función $f : n \rightarrow A$ y se denota

$$\langle a_i \mid i < n \rangle.$$

La **secuencia vacía** es la única secuencia de longitud cero. Denotamos $\text{Seq}(A) = \bigcup_{n \in \mathbb{N}} A^n$, el conjunto de todas las secuencias finitas de elementos de A .

Una **secuencia infinita** (o simplemente **sucesión**) con valores en A es una función $f : \mathbb{N} \rightarrow A$, denotada por

$$\langle a_i \mid i \in \mathbb{N} \rangle.$$

donde $a_n = f(n)$ para cada $n \in \mathbb{N}$.

El siguiente resultado formaliza el hecho de que proporcionar un dato inicial y una regla recursiva determina una única sucesión.

Teorema 4 (Recursión) Sea A un conjunto, sea $a \in A$ y sea $g : A \times \mathbb{N} \rightarrow A$ una función. Entonces existe una única sucesión $f : \mathbb{N} \rightarrow A$ tal que

$$f(0) = a, \quad f(n+1) = g(f(n), n) \quad \forall n \in \mathbb{N}.$$

En otras palabras, el dato del valor inicial a y de la regla recursiva g determina de manera unívoca una sucesión $\langle f(n) \mid n \in \mathbb{N} \rangle$ que cumple las ecuaciones anteriores, y recíprocamente, toda sucesión definida por esas ecuaciones procede de algún par (a, g) .

Demostración. La demostración se basa en considerar todas las funciones parciales h cuyo dominio es un segmento inicial de $\mathbb{N}$ y que satisfacen las ecuaciones de recurrencia en dicho dominio, y en mostrar que la unión de todas ellas es una función $f : \mathbb{N} \rightarrow A$ bien definida que cumple las condiciones del enunciado y es necesariamente única (véase [11, Cap. 3]). $\square$

Aplicaremos este principio para fijar rigurosamente la suma, el producto y la potenciación sobre los números naturales. Estas operaciones se definen por recursión sobre el segundo argumento.

Definición 30 (Operadores por recursión) Para $m, n \in \mathbb{N}$, definimos:

1. Suma:

$$m + 0 := m, \quad m + (n + 1) := (m + n) + 1.$$

2. Producto:

$$m \cdot 0 := 0, \quad m \cdot (n + 1) := (m \cdot n) + m.$$

3. Potencia:

$$m^0 := 1, \quad m^{n+1} := m^n \cdot m.$$

Más precisamente, para cada $m \in \mathbb{N}$, las igualdades anteriores describen una sucesión

$$\langle m + n \mid n \in \mathbb{N} \rangle, \quad \langle m \cdot n \mid n \in \mathbb{N} \rangle, \quad \langle m^n \mid n \in \mathbb{N} \rangle$$

2.4. Conjuntos finitos, numerables y no numerables

definida recursivamente en n . El Teorema 4 garantiza que estas ecuaciones determinan de manera única funciones

$$+ : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}, \quad \cdot : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}, \quad (\cdot)^{(\cdot)} : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}.$$

A partir de las definiciones recursivas anteriores y por inducción sobre $\mathbb{N}$, se comprueba que la suma y el producto satisfacen las propiedades usuales: conmutatividad y asociatividad de $+$ y $\cdot$, distributividad del producto respecto de la suma y la regla de que $m \cdot n = 0$ con $m \neq 0$ implica $n = 0$.

Por todo ello, $(\mathbb{N}, 0, S)$ satisface el esquema axiomático de Peano: el sucesor S es inyectivo, 0 no es sucesor de ningún elemento y vale el principio de inducción (todo subconjunto de $\mathbb{N}$ que contiene a 0 y es estable por S coincide con $\mathbb{N}$). En consecuencia, las operaciones suma $(+)$ y producto $(\cdot)$ definidas por recursión hacen de $(\mathbb{N}, 0, 1, S, +, \cdot)$ un modelo de la aritmética usual de los números naturales, en el sentido de los axiomas de Peano (véase [14]).

2.4. Conjuntos finitos, numerables y no numerables

En esta sección formalizamos la idea intuitiva de **tamaño** de un conjunto. El punto de partida será la noción de **equipotencia**. A partir de ahí introduciremos un orden entre cardinalidades comparándolas mediante inyecciones, y utilizaremos el teorema de Cantor–Bernstein para pasar de desigualdades en ambos sentidos a biyecciones. Sobre esta base distinguiremos los **conjuntos finitos**, **conjuntos numerables** y cerraremos con el paso a **conjuntos no numerables**. De este modo, la sección ilustra el comportamiento de las distintas cardinalidades y prepara el terreno para la aritmética de números cardinales y la jerarquía de los números aleph desarrolladas más adelante.

2.4.1. Cardinalidad de conjuntos

La noción de **equipotencia** captura cuándo dos conjuntos tienen el mismo tamaño, es decir, existe una biyección entre ellos.

Definición 31 (Equipotencia) *Dos conjuntos A y B son **equipotentes** si existe una biyección $f : A \rightarrow B$. En tal caso, decimos que comparten **cardinalidad** y escribimos $|A| = |B|$.*

Esta noción captura con precisión la idea de “tener el mismo tamaño”: no importa contar, sino poder establecer una biyección que empareje los elementos sin sobrantes ni faltantes. En consecuencia, la equipotencia es una relación de equivalencia (reflexiva, simétrica y transitiva) que particiona el universo de conjuntos en clases de igual cardinalidad.

Definición 32 (Orden de cardinalidades) *Escribimos $|A| \leq |B|$ si existe una función inyectiva de A en B .*

Para comparar tamaños permitiendo desigualdad, trabajaremos con inyecciones; la clave técnica es que inyecciones en ambos sentidos fuerzan la igualdad de tamaños.

Teorema 5 (Cantor–Bernstein) Sean A y B conjuntos. Si $|A| \leq |B|$ y $|B| \leq |A|$, entonces $|A| = |B|$.

Una demostración de este teorema puede encontrarse en [11, pág. 66].

Además, si $A_1 \subseteq B \subseteq A$ y A_1 es equipotente a A , entonces cualquier conjunto intermedio B tiene la misma cardinalidad que A .

En esta memoria, hablaremos de **números cardinales** tales que para cada conjunto A existe un objeto (su **cardinal**, que denotamos $|A|$) tal que dos conjuntos tienen el mismo cardinal si y sólo si son equipotentes.

2.4.2. Conjuntos finitos, numerables y no numerables

De manera intuitiva, diremos que un conjunto es finito cuando sus elementos se pueden “contar” con un número natural: existe un $n \in \mathbb{N}$ tal que el conjunto tiene exactamente n elementos.

Definición 33 (Conjunto finito) Un conjunto A es **finito** si existe $n \in \mathbb{N}$ con $A \sim n$. Si no es finito, es **infinito**.

Proposición 6 (Subconjuntos de un conjunto finito) Si A es finito y $B \subseteq A$, entonces B es finito y se cumple $|B| \leq |A|$.

Demostración. Sea $|A| = n$ para algún $n \in \mathbb{N}$. Demostraremos la afirmación por inducción sobre n . Si $n = 0$, entonces $A = \emptyset$ y su único subconjunto es $\emptyset$, que es finito. Supongamos ahora que el resultado es cierto para todos los conjuntos de cardinal n y consideremos un conjunto A con $|A| = n + 1$. Fijamos un elemento $a \in A$ y escribimos $A' := A \setminus \{a\}$, de modo que $|A'| = n$. Sea $B \subseteq A$. Distinguimos dos casos:

- Si $a \notin B$, entonces $B \subseteq A'$ y, por hipótesis de inducción, B es finito.
- Si $a \in B$, ponemos $B' := B \setminus \{a\}$, de manera que $B' \subseteq A'$. De nuevo, por hipótesis de inducción, B' es finito, es decir, $|B'| = m$ para algún $m \in \mathbb{N}$. Entonces $B = B' \cup \{a\}$ tiene exactamente $m + 1$ elementos, luego también es finito.

En ambos casos, B es finito. Además, en todos los pasos el número de elementos de B nunca excede $n + 1 = |A|$, así que $|B| \leq |A|$. $\square$

Proposición 7 (Uniones finitas) Si A y B son finitos, entonces $A \cup B$ es finito. Además, si $A \cap B = \emptyset$, entonces $|A \cup B| = |A| + |B|$.

Demostración. Sean $|A| = n$ y $|B| = m$ con $n, m \in \mathbb{N}$. Como A y B son finitos, podemos escribirlos como $A = \{a_0, \dots, a_{n-1}\}$, $B = \{b_0, \dots, b_{m-1}\}$. Consideremos ahora la sucesión finita $z = \{a_0, \dots, a_{n-1}, b_0, \dots, b_{m-1}\}$, de longitud $n + m$. Por construcción, todo elemento de $A \cup B$ aparece en la sucesión finita z , luego $A \cup B$ es imagen de un conjunto finito y, por tanto, es finito. Además, esto muestra que $|A \cup B| \leq n + m$.

Si, además, $A \cap B = \emptyset$, ningún elemento de A coincide con uno de B , de modo que todos los elementos de z son distintos. En ese caso, la sucesión (z_i) es biyectiva sobre $A \cup B$, y se obtiene $|A \cup B| = n + m = |A| + |B|$. $\square$

2.4. Conjuntos finitos, numerables y no numerables

En particular, si A es infinito, su cardinal excede cualquier número natural, es decir, para todo $n \in \mathbb{N}$, $|A| > n$.

Teorema 6 ($\mathbb{N}$ es infinito) *El conjunto $\mathbb{N}$ no es finito.*

Demostración. Supongamos, por absurdo, que $\mathbb{N}$ es finito. Entonces el conjunto $\mathbb{N}$ tendría un máximo m . Sin embargo, por la definición de sucesor, $m + 1$ es un número natural, es decir, $m + 1 \in \mathbb{N}$, y además $m < m + 1$. Esto contradice el hecho de que m sea máximo de $\mathbb{N}$. Por tanto, $\mathbb{N}$ no puede ser finito. $\square$

Llamaremos **numerable** a todo conjunto equipotente con $\mathbb{N}$ y **a lo sumo numerable** a todo aquel que se inyecta en $\mathbb{N}$.

Definición 34 (Numerable) *Un conjunto A es **numerable** si $|A| = |\mathbb{N}|$. Se dice que es **a lo sumo numerable** si $|A| \leq |\mathbb{N}|$.*

La clase de los conjuntos numerables es notablemente estable, en el sentido recogido en la siguiente proposición.

Proposición 8 (Uniones de conjuntos numerables) *La unión de dos conjuntos numerables es numerable.*

Demostración. Sean $A = \{a_n \mid n \in \mathbb{N}\}$ y $B = \{b_n \mid n \in \mathbb{N}\}$ enumeraciones de A y B . Construimos una sucesión $(c_k)_{k \in \mathbb{N}}$ definiendo $c_{2n} = a_n$, $c_{2n+1} = b_n$, ($n \in \mathbb{N}$). Entonces $\{c_k \mid k \in \mathbb{N}\}$ contiene a todos los elementos de $A \cup B$, luego $A \cup B$ es a lo sumo numerable. Como $A \cup B$ es infinito (contiene a A), eliminando las posibles repeticiones de la sucesión (c_k) obtenemos una nueva sucesión $(d_k)_{k \in \mathbb{N}}$ de elementos distintos cuya imagen es exactamente $A \cup B$. La aplicación $k \mapsto d_k$ es una biyección $\mathbb{N} \rightarrow A \cup B$, por lo que $A \cup B$ es numerable. $\square$

Teorema 7 ($\mathbb{Z}$ y $\mathbb{Q}$ numerables) *El conjunto de los enteros $\mathbb{Z}$ y el conjunto de los racionales $\mathbb{Q}$ son numerables.*

Demostración.

(1) $\mathbb{Z}$ numerable. Podemos escribir $\mathbb{Z} = \{0, 1, 2, 3, \dots\} \cup \{-1, -2, -3, \dots\}$. El subconjunto $\{0, 1, 2, \dots\}$ es equipotente a $\mathbb{N}$ mediante la identidad, y

$$f : \mathbb{N} \rightarrow \{-1, -2, -3, \dots\}, \quad f(n) = -(n + 1),$$

es una biyección, de modo que $\{-1, -2, -3, \dots\}$ también es numerable. Como la unión de dos conjuntos numerables es numerable, se concluye que $\mathbb{Z}$ es numerable.

(2) $\mathbb{Q}$ numerable. Véase un ejemplo de demostración en el siguiente Capítulo 3, concretamente en la Subsección 3.2.2. $\square$

Definición 35 (Aleph-cero) *Se introduce el símbolo $\aleph_0$ (aleph-cero) para denotar la cardinalidad de cualquier conjunto numerable. Así, $|\mathbb{N}| = |\mathbb{Z}| = |\mathbb{Q}| = \aleph_0$.*

En la Subsección 2.5.1 retomaremos estas operaciones a nivel cardinal y, en particular, las identidades básicas para $\aleph_0$.

El panorama del infinito no se agota en lo numerable: existen cardinalidades estrictamente mayores que $\aleph_0$. El ejemplo paradigmático es el conjunto de los números reales, que es no numerable.

Teorema 8 (No numerabilidad de $\mathbb{R}$ (Cantor)) *El conjunto $\mathbb{R}$ de los números reales es no numerable.*

Una demostración especialmente ilustrativa de este resultado, basada en el método diagonal de Cantor, puede encontrarse más adelante en el capítulo de paradojas (véase Subsección 3.2.2).

En la siguiente sección, el Teorema 10 mostrará el resultado general $|\mathcal{P}(X)| > |X|$; en particular, $|\mathcal{P}(\mathbb{N})| > |\mathbb{N}|$. Por su parte, remitimos a la Subsección 2.5.2 para la identificación de la cardinalidad del continuo y sus consecuencias.

2.5. Números cardinales

2.5.1. Aritmética cardinal

Trabajaremos con operaciones entre cardinales que reflejan, a nivel de conjuntos, la unión disjunta, el producto cartesiano y el conjunto de funciones. Usaremos que dos conjuntos son equipotentes si existe una biyección entre ellos, y escribiremos $|A| = \kappa$ cuando el cardinal de A sea κ .

Definición 36 (Operaciones aritméticas de cardinales) *Dados conjuntos disjuntos A, B con $|A| = \kappa$ y $|B| = \lambda$, definimos*

$$\kappa + \lambda := |A \cup B| \quad (A \cap B = \emptyset), \quad \kappa \cdot \lambda := |A \times B|, \quad \kappa^\lambda := |A^B|,$$

donde A^B es el conjunto de funciones $B \rightarrow A$.

Estas operaciones no dependen de los representantes concretos elegidos: si $|A| = |A'|$ y $|B| = |B'|$, entonces

$$|A \cup B| = |A' \cup B'|, \quad |A \times B| = |A' \times B'|, \quad |A^B| = |A'^{B'}|.$$

La prueba combina biyecciones dadas $A \cong A'$, $B \cong B'$ con copias disjuntas, productos de biyecciones y transporte por composición en A^B .

Proposición 9 (Leyes básicas) *La suma, producto y potencia de cardinales satisfacen las siguientes propiedades:*

1. **Conmutatividad:** $\kappa + \lambda = \lambda + \kappa$ y $\kappa \cdot \lambda = \lambda \cdot \kappa$.
2. **Asociatividad:** $\kappa + (\lambda + \mu) = (\kappa + \lambda) + \mu$ y $\kappa \cdot (\lambda \cdot \mu) = (\kappa \cdot \lambda) \cdot \mu$.
3. **Distributividad:** $\kappa \cdot (\lambda + \mu) = \kappa \cdot \lambda + \kappa \cdot \mu$.
4. **Monotonía:** si $\kappa_1 \leq \kappa_2$ y $\lambda_1 \leq \lambda_2$, entonces $\kappa_1 + \lambda_1 \leq \kappa_2 + \lambda_2$ y $\kappa_1 \cdot \lambda_1 \leq \kappa_2 \cdot \lambda_2$.

Para probarlas, bastaría con traducir cada operación con cardinales a la operación correspondiente con conjuntos y construir biyecciones explícitas que muestren la conmutatividad, asociatividad, distributividad y monotonía de estas operaciones.

Teorema 9 (Leyes de exponentes) *La suma, producto y potencia de cardinales satisfacen las siguientes propiedades:*

1. $\kappa^{\lambda+\mu} = \kappa^\lambda \cdot \kappa^\mu$.
2. $(\kappa^\lambda)^\mu = \kappa^{\lambda \cdot \mu}$.
3. $(\kappa \cdot \lambda)^\mu = \kappa^\mu \cdot \lambda^\mu$.

Demostración. Para esta demostración se ha de interpretar κ^λ como el cardinal del conjunto de funciones cuyo dominio tiene cardinal λ y codominio tiene cardinal κ . Con esta lectura: (1) corresponde a descomponer el dominio en una unión disjunta de dos partes (de cardinales λ y μ), (2) a identificar una familia de funciones indexada por un conjunto de tamaño μ con una sola función definida en el producto de dominios (de tamaño $\lambda \cdot \mu$), y (3) a elegir, para cada punto del dominio, un par formado por un valor en un conjunto de tamaño κ y otro en uno de tamaño λ , lo que equivale a elegir de manera independiente cada componente. Una demostración detallada puede encontrarse en [11, págs. 95–96]. $\square$

El crecimiento de potencias de conjuntos está controlado por el resultado clásico de Cantor:

Teorema 10 (Conjunto potencia: versión general de Cantor) *Para todo conjunto X , se cumple $|\mathcal{P}(X)| > |X|$.*

Demostración. Consideremos la aplicación $i : X \rightarrow \mathcal{P}(X)$, $i(x) = \{x\}$. Si $i(x) = i(y)$, entonces $\{x\} = \{y\}$ y, por extensionalidad, $x = y$, de modo que i es inyectiva y obtenemos $|X| \leq |\mathcal{P}(X)|$. Nos falta ver que no existe ninguna aplicación sobreyectiva de X sobre $\mathcal{P}(X)$. Sea $f : X \rightarrow \mathcal{P}(X)$ arbitraria y definamos $S := \{x \in X \mid x \notin f(x)\}$. Afirmamos que S no está en la imagen de f . En efecto, si existiera $z \in X$ tal que $f(z) = S$, por la definición de S tendríamos $z \in S$ si y sólo si $z \notin f(z) = S$, lo cual es imposible. Por tanto ninguna aplicación $X \rightarrow \mathcal{P}(X)$ es sobreyectiva. En conclusión, se cumple $|X| \leq |\mathcal{P}(X)|$ y, además, $|X| \neq |\mathcal{P}(X)|$, por lo que necesariamente $|X| < |\mathcal{P}(X)|$. $\square$

Interpretando subconjuntos como funciones características $X \rightarrow \{0, 1\}$ se obtiene además

$$|\mathcal{P}(X)| = 2^{|X|},$$

por lo que el Teorema 10 afirma que para todo cardinal κ se tiene $\kappa < 2^\kappa$. En particular, no existe un cardinal máximo.

2.5.2. La cardinalidad del continuo

Como caso testigo de la aritmética cardinal, conviene entender bien cómo se comporta el cardinal $\aleph_0$ y cómo se relaciona con el cardinal de los números reales.

Proposición 10 (Aritmética con $\aleph_0$) *Para todo $n \in \mathbb{N}$ con $n > 0$:*

1. $n + \aleph_0 = \aleph_0$.
2. $n \cdot \aleph_0 = \aleph_0$.

$$3. \aleph_0^n = \aleph_0.$$

Estas igualdades se obtienen combinando desigualdades cardinales evidentes, el hecho de que la unión y el producto de conjuntos numerables siguen siendo numerables, y una aplicación del teorema de Cantor–Bernstein (Teorema 5). Además, estas biyecciones admiten una interpretación muy intuitiva en términos del Hotel infinito de Hilbert, que se describirá con detalle en la Subsección 3.2.4.

El siguiente resultado fija el cardinal del conjunto de los números reales y sirve de referencia para el resto del capítulo.

Teorema 11 (Cardinalidad del continuo) *Se cumple $|\mathbb{R}| = 2^{\aleph_0}$.*

Demostración. Dividiremos la prueba en varios pasos.

1) Identificación de $2^{\aleph_0}$ con $\mathcal{P}(\mathbb{N})$.

Sea $\{0, 1\}^{\mathbb{N}}$ el conjunto de todas las sucesiones binarias $(a_n)_{n \in \mathbb{N}}$. Por definición de potencia de cardinales,

$$|\{0, 1\}^{\mathbb{N}}| = 2^{\aleph_0}.$$

Por otra parte, a cada subconjunto $A \subseteq \mathbb{N}$ le asociamos su función característica $\chi_A : \mathbb{N} \rightarrow \{0, 1\}$,

$$\chi_A(n) = \begin{cases} 1 & \text{si } n \in A, \\ 0 & \text{si } n \notin A. \end{cases}$$

La aplicación $A \rightarrow \chi_A$ induce una biyección entre $\mathcal{P}(\mathbb{N})$ y $\{0, 1\}^{\mathbb{N}}$ (a cada sucesión binaria le corresponde un único subconjunto de $\mathbb{N}$, y recíprocamente). En consecuencia,

$$|\mathcal{P}(\mathbb{N})| = |\{0, 1\}^{\mathbb{N}}| = 2^{\aleph_0}.$$

2) Una inyección $\mathcal{P}(\mathbb{N}) \rightarrow (0, 1)$.

Definimos una aplicación $\Phi : \mathcal{P}(\mathbb{N}) \rightarrow (0, 1)$ asociando a cada subconjunto $A \subseteq \mathbb{N}$ el número real

$$\Phi(A) := \sum_{n \in A} 10^{-n}.$$

Esta serie converge y define un número en $(0, 1)$, cuya expansión decimal comienza por $0.d_1d_2d_3\dots$. Dos subconjuntos distintos $A \neq B$ producen secuencias de dígitos distintas, y por tanto números distintos. Así, Φ es inyectiva, y se obtiene

$$|\mathcal{P}(\mathbb{N})| \leq |(0, 1)|.$$

3) Una inyección $(0, 1) \rightarrow \mathcal{P}(\mathbb{N})$.

Todo número real $x \in (0, 1)$ admite al menos una expansión decimal

$$x = 0.a_1a_2a_3\dots, \quad a_n \in \{0, 1, \dots, 9\}.$$

La única posible ambigüedad proviene de las expansiones que terminan en una cola infinita de 9. Por convenio, para cada x elegimos la expansión que no termina en una cola infinita de 9; de este modo, cada $x \in (0, 1)$ queda asociado a una única sucesión de dígitos $(a_n)_{n \in \mathbb{N}}$.

Codificamos ahora esos dígitos como subconjuntos de $\mathbb{N}$. Por ejemplo, fijemos una biyección fija $\pi : \mathbb{N} \times \{0, 1, \dots, 9\} \rightarrow \mathbb{N}$, y para cada $x \in (0, 1)$ con expansión $0.a_1a_2a_3\dots$ definimos

$$E(x) := \{ \pi(n, k) \mid n \in \mathbb{N}, 0 \leq k < a_n \} \subseteq \mathbb{N}.$$

Si $x \neq y$, sus expansiones decimales canónicas difieren en algún dígito, digamos $a_m \neq b_m$; esto implica $E(x) \neq E(y)$, luego la aplicación $\Psi : (0, 1) \rightarrow \mathcal{P}(\mathbb{N})$, $\Psi(x) := E(x)$, es inyectiva. Por tanto,

$$|(0, 1)| \leq |\mathcal{P}(\mathbb{N})|.$$

4) Conclusión.

De los pasos 2) y 3) obtenemos

$$|\mathcal{P}(\mathbb{N})| \leq |(0, 1)| \leq |\mathcal{P}(\mathbb{N})|,$$

y por el teorema de Cantor–Bernstein (Teorema 5) se concluye

$$|(0, 1)| = |\mathcal{P}(\mathbb{N})|.$$

Finalmente, ya sabemos que $(0, 1)$ y $\mathbb{R}$ son equipotentes, de modo que

$$|\mathbb{R}| = |(0, 1)| = |\mathcal{P}(\mathbb{N})| = 2^{\aleph_0}.$$

□

En adelante llamaremos **cardinalidad del continuo** al cardinal $2^{\aleph_0}$ y lo denotaremos por $\mathfrak{c}$. Así, $|\mathbb{R}| = \mathfrak{c}$ y, de manera equivalente, $|2^{\mathbb{N}}| = \mathfrak{c}$. Con esta notación, las operaciones finitas no modifican el tamaño del continuo:

Proposición 11 (Operaciones con el continuo) Para todo $n \in \mathbb{N}$, $n > 0$:

1. $n + \mathfrak{c} = \mathfrak{c}$.
2. $n \cdot \mathfrak{c} = \mathfrak{c}$.
3. $\mathfrak{c}^n = \mathfrak{c}$.

Demostración. Recordemos que $\mathfrak{c} = 2^{\aleph_0}$ y que $\aleph_0 < \mathfrak{c}$ (Teoremas 8 y 11). Usaremos además que $\aleph_0 + \aleph_0 = \aleph_0$ y $\aleph_0 \cdot \aleph_0 = \aleph_0$ (Proposición 10) y las leyes de exponentes (Teorema 9).

(1) Para la suma, encajamos $n + \mathfrak{c}$ entre dos copias de $\mathfrak{c}$:

$$\mathfrak{c} \leq n + \mathfrak{c} \leq \aleph_0 + \mathfrak{c} \leq \mathfrak{c} + \mathfrak{c} = 2 \cdot \mathfrak{c} = \mathfrak{c}.$$

Por el teorema de Cantor–Bernstein (Teorema 5) se obtiene $n + \mathfrak{c} = \mathfrak{c}$.

(2) Para el producto se procede de forma análoga:

$$\mathfrak{c} \leq n \cdot \mathfrak{c} \leq \aleph_0 \cdot \mathfrak{c} \leq \mathfrak{c} \cdot \mathfrak{c} = \mathfrak{c}^2 = \mathfrak{c},$$

de donde, de nuevo por Cantor–Bernstein, $n \cdot \mathfrak{c} = \mathfrak{c}$.

Capítulo 2. Teoría de conjuntos: nociones y conceptos básicos

(3) Para la potencia, usamos que $n > 0$:

$$\mathfrak{c} \leq \mathfrak{c}^n \leq \mathfrak{c}^{\aleph_0} = \mathfrak{c}.$$

Aplicando una vez más Cantor–Bernstein concluimos $\mathfrak{c}^n = \mathfrak{c}$. □

En particular, muchas construcciones habituales comparten la cardinalidad del continuo. Por ejemplo, para todo $n \in \mathbb{N}$ con $n > 0$ el espacio euclídeo $\mathbb{R}^n$ tiene cardinalidad $\mathfrak{c}$; en el caso $n = 2$ esto dice que el plano real $\mathbb{R}^2$ tiene cardinalidad $\mathfrak{c}$. Como el cuerpo de los complejos $\mathbb{C}$ se identifica canónicamente con $\mathbb{R}^2$ mediante $x + iy \iff (x, y)$, también se cumple $|\mathbb{C}| = \mathfrak{c}$. Además, el conjunto de todas las sucesiones infinitas de reales se puede ver como $\mathbb{R}^{\mathbb{N}}$, y de igual modo las sucesiones de naturales forman $\mathbb{N}^{\mathbb{N}}$; en ambos casos la cardinalidad vuelve a ser $\mathfrak{c}$.

Además, eliminar un subconjunto numerable no reduce el tamaño del continuo: si A es numerable y B satisface $|B| = \mathfrak{c}$, entonces $|B \setminus A| = \mathfrak{c}$. En efecto, B se descompone como unión disjunta $B = A \cup (B \setminus A)$, de modo que

$$|B| = |A| + |B \setminus A| = \aleph_0 + |B \setminus A|.$$

Por la Proposición 11, sumar un cardinal numerable no aumenta el continuo, así que necesariamente $|B \setminus A| = \mathfrak{c}$.

Como consecuencia inmediata, también tienen cardinalidad $\mathfrak{c}$ los números irracionales, el conjunto de todos los subconjuntos infinitos de $\mathbb{N}$ y el de las biyecciones $\mathbb{N} \rightarrow \mathbb{N}$, ya que en cada caso se obtienen a partir de un conjunto de cardinalidad $\mathfrak{c}$ eliminando únicamente un subconjunto numerable.

Finalmente, situamos la cuestión estructural que guiará varias referencias posteriores: ¿existen cardinales intermedios entre lo numerable y el continuo?

Definición 37 (Hipótesis del Continuo) *La **Hipótesis del Continuo** afirma que no existe un cardinal κ tal que*

$$\aleph_0 < \kappa < \mathfrak{c}.$$

En otras palabras, cualquier subconjunto infinito de $\mathbb{R}$ tiene cardinalidad numerable o bien la del continuo.

En presencia del Axioma de Elección (ZFC) todo conjunto puede ser bien ordenado, y cada cardinal infinito es equipotente con un **único ordinal inicial**. Estos ordinales iniciales se indexan como $\aleph_0, \aleph_1, \aleph_2, \dots, \aleph_\alpha, \dots$, de modo que $\aleph_1$ es, por definición, el menor cardinal estrictamente mayor que $\aleph_0$ (véase, por ejemplo, ([11, págs. 98–101])). En este contexto, la Hipótesis del Continuo es equivalente a la igualdad $\mathfrak{c} = 2^{\aleph_0} = \aleph_1$. En ausencia del Axioma de Elección la situación es más sutil: no todo conjunto es bien ordenable, el continuo puede no coincidir con ningún $\aleph_\alpha$, y la formulación habitual de CH deja de ser equivalente a la identidad $2^{\aleph_0} = \aleph_1$.

2.6. Números ordinales

2.6.1. Conjuntos bien ordenados

De manera informal, diremos que un orden es bueno cuando siempre existe un “punto de partida”: al considerar cualquier colección no vacía de elementos, hay uno que aparece antes que todos los demás según la relación $<$.

Definición 38 (Bien ordenado) Un par $(W, <)$ está **bien ordenado** si $<$ es un orden lineal en W y todo subconjunto no vacío de W tiene un mínimo.

Resulta útil fijar la noción de **segmento inicial** como el conjunto de todos los predecesores de un punto. La utilidad de los segmentos iniciales es que capturan la idea de “todo lo que viene antes” de un punto del orden; serán la pieza clave al comparar dos conjuntos bien ordenados.

Definición 39 (Segmento inicial) Sea $(W, <)$ bien ordenado. Un subconjunto $S \subseteq W$ es un **segmento inicial** si existe $a \in W$ tal que

$$S = \{x \in W : x < a\}.$$

Al punto a lo llamaremos el **corte** que determina S .

Diremos que dos conjuntos ordenados $(A, \leq_A)$ y $(B, \leq_B)$ son **isomorfos** si existe una biyección $f : A \rightarrow B$ tal que, para cualesquiera $x, y \in A$,

$$x \leq_A y \iff f(x) \leq_B f(y).$$

En ese caso, f se llama un **isomorfismo de órdenes** entre $(A, \leq_A)$ y $(B, \leq_B)$.

Con esta noción, cualesquiera dos conjuntos bien ordenados se comparan de forma canónica: o bien son isomorfos, o bien uno es isomorfo a un segmento inicial del otro. La siguiente proposición precisa este principio de comparación.

Proposición 12 (Comparación de bien ordenados) Si $(W_1, <_1)$ y $(W_2, <_2)$ son conjuntos bien ordenados, entonces exactamente una de las siguientes alternativas ocurre:

1. W_1 y W_2 son isomorfos;
2. W_1 es isomorfo a un segmento inicial de W_2 ;
3. W_2 es isomorfo a un segmento inicial de W_1 .

En todos los casos, el isomorfismo es único.

La demostración de este resultado excede el alcance de este trabajo, por lo que la omitimos. Puede consultarse en [11, págs. 105–106].

2.6.2. Números ordinales

Recordemos que $(\mathbb{N}, <)$ es un **conjunto bien ordenado** (Teorema 3); los **ordinales** extienden esta situación a órdenes transfinitos. Diremos que un conjunto T es **transitivo** si todo elemento de T es a su vez subconjunto de T . Con esta terminología, definimos los números ordinales como sigue.

Capítulo 2. Teoría de conjuntos: nociones y conceptos básicos

Definición 40 (Ordinal) Un conjunto α es un **ordinal** si es transitivo y está bien ordenado con la relación de pertenencia $\in$.

Proposición 13 (Sucesor ordinal) Si α es ordinal, entonces su sucesor $S(\alpha)$ también es ordinal.

Demostración. Por definición, $S(\alpha) = \alpha \cup \{\alpha\}$ es un conjunto transitivo. Además, $\alpha \cup \{\alpha\}$ está bien ordenado por $\in$, siendo α su mayor elemento, y $\alpha \subset \alpha \cup \{\alpha\}$ el segmento inicial determinado por α . En consecuencia, $S(\alpha)$ es un número ordinal. $\square$

Los naturales se insertan de manera canónica en la jerarquía ordinal.

Proposición 14 (Naturales como ordinales) Todo número natural es un ordinal.

Demostración. Procedemos por inducción sobre $n \in \mathbb{N}$.

Para $n = 0$, recordemos que $0 := \emptyset$. El conjunto vacío es transitivo y $(\emptyset, \in)$ está bien ordenado: cualquier subconjunto no vacío de $\emptyset$ es imposible, de modo que la condición de existir un mínimo se satisface de forma trivial. Por tanto, 0 es un ordinal.

Supongamos ahora que n es un ordinal. Entonces su sucesor $n^* := n \cup \{n\}$ es también un ordinal, por la Proposición 13. Por la construcción usual de los naturales, $n^* = n + 1$. De este modo, siempre que n sea un ordinal, también lo es $n + 1$, y por inducción tenemos que todo natural $n \in \mathbb{N}$ es un ordinal. $\square$

Como consecuencia, $\mathbb{N}$ es un conjunto de ordinales que contiene a 0 y es cerrado por sucesor. Además, ya sabemos que es infinito y minimal entre los conjuntos inductivos (Proposición 5). En la teoría de ordinales se define ω precisamente como el menor ordinal infinito, es decir, como el menor conjunto de ordinales que contiene a 0 y es cerrado por sucesor.

El orden $<$ entre ordinales hereda las buenas propiedades de los conjuntos bien ordenados. En particular, dados ordinales α y β , siempre se cumple exactamente una de las tres posibilidades $\alpha < \beta$, $\alpha = \beta$, o bien $\beta < \alpha$, de modo que el orden es lineal. Además, todo subconjunto no vacío de ordinales tiene mínimo respecto de $<$, en este sentido, la propia clase de los ordinales está bien ordenada por $<$.

Una consecuencia importante es que no existe un conjunto que contenga a todos los ordinales. En efecto, supongamos que existe un conjunto X cuyos elementos son todos los ordinales. Como todo ordinal es transitivo, la unión

$$\gamma := \bigcup X$$

es también transitiva, y $(\gamma, \in)$ está bien ordenado (hereda el buen orden de los elementos de X). Por tanto, γ es un ordinal. Pero entonces, al ser γ un ordinal, debería pertenecer a X , lo que implicaría $\gamma \in \gamma$, en contradicción con la Proposición 15. Concluimos que no existe un conjunto de todos los ordinales.

Proposición 15 (Elementariedad) Si α es ordinal, entonces $\alpha \notin \alpha$; además, todo elemento de un ordinal es ordinal y, si $\alpha \subset \beta$, entonces $\alpha < \beta$.

Demostración. Como el orden $<$ entre ordinales es, por construcción, un orden estricto bien fundado, en particular es irreflexivo; por tanto, para todo ordinal α se tiene $\alpha \not< \alpha$.

Sea ahora $\beta < \alpha$. Demostraremos que β es un ordinal. Como α es transitivo y $\beta < \alpha$, si $u < v < \beta$, se tiene $u, v < \alpha$, y por tanto $u < \beta$, ya que $<_\alpha$ es una relación de buen orden. Así, β es transitivo. Además, dado que $\beta \subseteq \alpha$, la relación $<_\beta$ es la restricción de $<_\alpha$ a β , y puesto que $<_\alpha$ es un buen orden, también lo es $<_\beta$. En consecuencia, β es un ordinal.

Por último, probemos que si $\alpha < \beta$, entonces $\alpha \in \beta$. Dado que $\beta \setminus \alpha$ es un subconjunto no vacío de β , el buen orden de $<_\beta$ garantiza que posee un menor elemento γ . Por construcción, $\gamma \not< \alpha$ y $\gamma \subseteq \alpha \cup \{\gamma\} \subseteq \beta$. Supongamos, por absurdo, que $\alpha \not< \beta$. Entonces α y γ serían elementos de β y, por tanto, o bien $\gamma < \alpha$ o bien $\alpha < \gamma$. Pero la primera opción contradice $\gamma \not< \alpha$, y la segunda implica que $\alpha < \beta$, como queríamos. $\square$

Todas estas propiedades elementales permiten identificar exactamente los ordinales finitos con los números naturales.

Proposición 16 (Ordinales finitos) *Los números naturales son exactamente los números ordinales finitos.*

Demostración. Por la Proposición 14, todo $n \in \omega = \mathbb{N}$ es un ordinal y, además, finito. Veamos la dirección recíproca.

Sea α un ordinal finito. Supongamos, buscando una contradicción, que no se cumple $\alpha < \omega$. Como α y ω son ordinales, por la comparación de ordinales sólo pueden darse tres casos: $\alpha < \omega$, $\alpha = \omega$, $\omega < \alpha$. Nuestra hipótesis descarta el primero. Además, α es finito mientras que ω es infinito, así que tampoco puede ocurrir $\alpha = \omega$. Por tanto, debe cumplirse $\omega < \alpha$.

Por la definición del orden entre ordinales, de $\omega < \alpha$ se sigue $\omega \in \alpha$. Como todo ordinal es transitivo, de $\omega \in \alpha$ se deduce $\omega \subseteq \alpha$, de modo que α contiene un subconjunto infinito. Esto contradice que α sea finito.

En consecuencia, nuestra suposición inicial es imposible y debe ocurrir $\alpha < \omega$, equivalentemente, $\alpha \in \omega$. Es decir, α es un número natural. $\square$

Teorema 12 (Representación por ordinal) *Todo conjunto bien ordenado es isomorfo a un único ordinal, llamado su **tipo de orden**.*

La demostración de este resultado se aleja del alcance de este trabajo, por lo que la omitimos. Obsérvese que el teorema extiende al caso transfinito el fenómeno que ya hemos probado para los casos finitos: todo conjunto bien ordenado finito es isomorfo a un único número natural (visto como ordinal). En esta memoria sólo se demuestra en detalle la versión finita; una prueba completa del caso general puede consultarse, por ejemplo, en [11, págs. 111–112].

Ahora, recordemos que identificamos ω con $\mathbb{N}$ (Definición 26) y que el orden natural está bien fundado (Teorema 3). El principio de **inducción transfinita** extiende la inducción sobre ω al universo de los ordinales: para demostrar una

propiedad $P(\alpha)$ para todo ordinal α , basta suponerla heredada de todos los ordinales menores.

Teorema 13 (Inducción transfinita) *Sea $P(\alpha)$ una propiedad definida para los ordinales. Supongamos que para cada ordinal α se cumple lo siguiente: Si $P(\beta)$ es verdadera para todo ordinal $\beta < \alpha$, entonces $P(\alpha)$ también es verdadera. Bajo esta hipótesis, $P(\alpha)$ es verdadera para todo ordinal α .*

Demostración. Supongamos, por absurdo, que existe algún ordinal γ que no satisface P , es decir, que $P(\gamma)$ es falsa, y consideremos el conjunto $S := \{\beta \leq \gamma \mid P(\beta) \text{ es falsa}\}$. Por hipótesis, S es no vacío; además, $S \subseteq \gamma + 1$ y, como todo segmento inicial de ordinales está bien ordenado por $<$, el conjunto S tiene un menor elemento, que denotamos por α . Entonces, por la minimalidad de α en S , para todo $\beta < \alpha$ se cumple $P(\beta)$. Aplicando ahora la hipótesis de inducción transfinita al ordinal α , concluimos que también debe cumplirse $P(\alpha)$, lo que contradice la definición de S , mostrando que S debe ser vacío, es decir, no hay ordinales que fallen la propiedad P . Por tanto, $P(\alpha)$ es verdadera para todo ordinal α . $\square$.

En la práctica se utiliza con frecuencia la siguiente versión “por casos” de la inducción transfinita: si se verifica $P(0)$, la herencia al sucesor $P(\alpha) \Rightarrow P(\alpha + 1)$ y, para ordinales límite α , la validez de $P(\beta)$ para todo $\beta < \alpha$ implica $P(\alpha)$, entonces $P(\alpha)$ es verdadera para todo ordinal α .

Teorema 14 (Recursión transfinita) *Sea g una regla de asignación que, a cada función h con dominio un ordinal α (es decir, $h : \alpha \rightarrow U$ para algún universo fijo de conjuntos U), le asocia un objeto $g(h)$. Entonces existe una única función f cuyo dominio es la clase de los ordinales tal que $f(\alpha) = g(f \upharpoonright \alpha)$, para todo ordinal α .*

Para la demostración de este teorema, se construye f por inducción transfinita: para cada ordinal α se supone ya definida la restricción $f \upharpoonright \alpha$ y se fija $f(\alpha) = g(f \upharpoonright \alpha)$. La inducción transfinita garantiza que estas definiciones parciales son compatibles y dan lugar a una única función f definida en todos los ordinales.

2.6.3. Aritmética ordinal

Al igual que en la aritmética cardinal (Subsección 2.5.1), ahora definimos la suma, el producto y la exponenciación de ordinales por **recursión** en el segundo argumento. Sin embargo, aquí las operaciones codifican tipos de orden y, en general, no son conmutativas.

Definición 41 (Suma, producto y potencia de ordinales) *Sea β un ordinal fijo. Para cualquier ordinal α y cualquier ordinal límite λ definimos, por inducción transfinita en el segundo argumento:*

1. Suma ordinal:

$$\beta + 0 = \beta, \quad \beta + (\alpha + 1) = (\beta + \alpha) + 1, \quad \beta + \lambda = \sup\{\beta + \gamma : \gamma < \lambda\}$$

2. Producto ordinal:

$$\beta \cdot 0 = 0, \quad \beta \cdot (\alpha + 1) = (\beta \cdot \alpha) + \beta, \quad \beta \cdot \lambda = \sup\{\beta \cdot \gamma : \gamma < \lambda\}$$

3. **Exponenciación ordinal:**

$$\beta^0 = 1, \quad \beta^{\alpha+1} = \beta^\alpha \cdot \beta, \quad \beta^\lambda = \sup\{\beta^\gamma : \gamma < \lambda\}$$

Proposición 17 (Leyes básicas) La suma y el producto ordinales son **asociativos** y **monótonos a la derecha**; en general **no conmutan**. Hay **cancelación a la izquierda**: si $\beta + \alpha_1 = \beta + \alpha_2$, entonces $\alpha_1 = \alpha_2$.

Demostración. La suma y el producto ordinales se han definido por recursión transfinita en el último argumento. A partir de esa definición, la asociatividad y la monotonía a la derecha se obtienen mediante inducción transfinita sobre dicho argumento, entendiendo $\beta + \alpha$ (resp. $\beta \cdot \alpha$) como el resultado de “pegar al final” copias ordenadas del tipo correspondiente. Para una demostración detallada remitimos a [11, págs. 120–122]. $\square$

La interpretación por **tipos de orden** conecta estas operaciones con construcciones sobre conjuntos bien ordenados. Trabajaremos con **copias disjuntas** de los conjuntos involucrados.

Definición 42 (Suma ordenada) Sean $(W_1, <_1)$ y $(W_2, <_2)$ bien ordenados y disjuntos. Definimos su **suma ordenada** como el conjunto $W_1 \cup W_2$ con la relación:

$$a < b \iff \begin{cases} a, b \in W_1 \text{ y } a <_1 b, & \text{o} \\ a, b \in W_2 \text{ y } a <_2 b, & \text{o} \\ a \in W_1, b \in W_2. & \end{cases}$$

Es decir, todos los elementos de W_1 preceden a todos los de W_2 , y dentro de cada bloque se mantiene el orden original.

Definición 43 (Órdenes lexicográfico y antilexicográfico) Sean $(\alpha, <)$ y $(\beta, <)$ ordinales, e interpretamos $\alpha \times \beta$ como producto cartesiano.

1. El **orden lexicográfico** compara primero la primera coordenada:

$$(a, b) <_{\text{lex}} (a', b') \iff a < a' \text{ o bien } (a = a' \text{ y } b < b').$$

2. El **orden antilexicográfico** compara primero la segunda coordenada:

$$(a, b) <_{\text{antilex}} (a', b') \iff b < b' \text{ o bien } (b = b' \text{ y } a < a').$$

Ambos definen buenos órdenes en $\alpha \times \beta$.

Proposición 18 (Interpretación por tipos de orden) Sea $(W_1, <_1)$ y $(W_2, <_2)$ bien ordenados de tipos α_1 y α_2 . Entonces:

1. La suma ordenada “ W_1 seguida de W_2 ” tiene tipo $\alpha_1 + \alpha_2$.
2. En $\alpha \times \beta$, la orden **antilexicográfica** tiene tipo $\alpha \cdot \beta$ y la **lexicográfica**, tipo $\beta \cdot \alpha$.

Capítulo 2. Teoría de conjuntos: nociones y conceptos básicos

No daremos la prueba en detalle. En ambos casos se construyen isomorfismos explícitos entre los órdenes considerados y los ordinales correspondientes (para la suma, una unión disjunta ordenada; para el producto, los órdenes lexicográfico y antilexicográfico sobre $\alpha \times \beta$). Además, se puede entender de una forma intuitiva en la que $\alpha_1 + \alpha_2$ representa “recorrer primero todos los elementos de W_1 y después todos los de W_2 ”, mientras que $\alpha \cdot \beta$ corresponde a “ β copias consecutivas del orden α ”. En particular, si identificamos ω con el orden $0 < 1 < 2 < \dots$, podemos describir explícitamente algunos ordinales básicos:

1. $\omega + 1$ es la sucesión

$$0 < 1 < 2 < \dots < \omega,$$

es decir, todos los naturales seguidos de un nuevo máximo.

2. $\omega + \omega = \omega \cdot 2$ es dos copias consecutivas de ω :

$$0 < 1 < 2 < \dots < \omega < \omega + 1 < \omega + 2 < \dots,$$

donde el “bloque” $\{\omega + n : n \in \omega\}$ viene después de todos los naturales.

3. $\omega^2 = \omega \cdot \omega$ puede verse como ω bloques consecutivos, cada uno ordenado como ω :

$$0 < 1 < 2 < \dots < \omega < \omega + 1 < \omega + 2 < \dots < \omega \cdot 2 < \omega \cdot 2 + 1 < \dots < \omega \cdot 3 < \dots,$$

donde para cada $n \in \omega$ el segmento $\{\omega \cdot n + k : k \in \omega\}$ es una copia de ω , y los bloques se suceden en el orden de $n = 0, 1, 2, \dots$.

Finalmente, todo ordinal positivo admite una **representación canónica** como suma finita de potencias decrecientes de ω con coeficientes naturales. Esto es la forma normal de Cantor:

Teorema 15 (Forma normal de Cantor) *Todo ordinal $\alpha > 0$ se escribe de manera única como*

$$\alpha = \omega^{\beta_1} \cdot k_1 + \omega^{\beta_2} \cdot k_2 + \dots + \omega^{\beta_n} \cdot k_n,$$

donde $n \geq 1$, $\beta_1 > \beta_2 > \dots > \beta_n$ son ordinales y $k_1, \dots, k_n \in \mathbb{N}$, $k_i > 0$.

Así, por ejemplo, $\omega = \omega^1 \cdot 1$, $\omega + 3 = \omega^1 \cdot 1 + \omega^0 \cdot 3$, $\omega \cdot 2 = \omega^1 \cdot 2$, y $\omega^2 + 5\omega + 7 = \omega^2 \cdot 1 + \omega^1 \cdot 5 + \omega^0 \cdot 7$. Ordinales más grandes, como ω^ω , también admiten una expresión de este tipo (en este caso, simplemente $\omega^\omega = \omega^\omega \cdot 1$). De este modo, muchas cuestiones sobre la aritmética ordinal pueden trasladarse al estudio de los exponentes y coeficientes que aparecen en su forma normal.

La demostración de este resultado se aleja del alcance de este trabajo, por lo que la omitimos. Puede consultarse, por ejemplo, en [11, págs. 124–125].

2.7. Cardinales y ordinales: números aleph

En esta sección usamos de forma esencial el **Axioma de Elección** (Axioma 10): bajo ZFC, todo cardinal infinito es equipotente con un **único número aleph** $\aleph_\alpha$, lo que nos permite indexar la escala de cardinales infinitos mediante ordinales.

2.7.1. Ordinales iniciales

Hasta ahora hemos identificamos cardinalidades como clases de equipotencia de conjuntos, y, en el caso bien ordenado, cada clase tiene un tipo de orden representado por un ordinal (Teorema 12). Ahora refinamos esta situación seleccionando, dentro de cada cardinalidad bien ordenable, un **primer ordinal** que la represente.

Definición 44 (Ordinal inicial) *Un ordinal α es **ordinal inicial** si no existe ningún $\beta < \alpha$ tal que $|\beta| = |\alpha|$.*

Es decir, α es el menor ordinal que tiene su cardinalidad. Entre los primeros ejemplos, todos los números naturales son ordinales iniciales. También ω lo es (no es equipotente con ningún natural); en cambio, ordinales como $\omega + 1$, $\omega + 2$, $\omega \cdot 2$ o ω^ω siguen siendo numerables, es decir, tienen la misma cardinalidad que ω , de modo que no son iniciales.

Por la Proposición 12, toda clase de equipotencia de conjuntos bien ordenados contiene a lo sumo un representante “más pequeño” (segmento inicial). De ahí se deduce que dentro de cada cardinalidad bien ordenable hay, como mucho, un ordinal inicial que la represente.

Teorema 16 (Representante inicial) *Si $(W, <)$ es un conjunto bien ordenado, entonces W es equipotente con un único ordinal inicial.*

Demostración. Por el Teorema 12, W es equipotente a algún ordinal α . Sea α_0 el menor ordinal equipotente con W . Entonces α_0 es inicial: si existiera $\beta < \alpha_0$ con $|\beta| = |\alpha_0|$, tendríamos $|W| = |\alpha_0| = |\beta|$, y β sería también equipotente con W , contradiciendo la minimalidad de α_0 .

Para la unicidad, si $\alpha_0 \neq \alpha_1$ son ordinales iniciales equipotentes con W , sin pérdida de generalidad $\alpha_0 < \alpha_1$. Pero entonces $|\alpha_0| = |\alpha_1|$, lo que contradice que α_1 sea inicial. $\square$

Esto motiva identificar, para conjuntos bien ordenables, su cardinalidad con dicho representante inicial. Si un conjunto W es bien ordenable, llamaremos **número cardinal** de W al único ordinal inicial equipotente con él, y lo denotaremos por $|W|$. En particular, $|\omega| = \aleph_0$ para todo conjunto numerable y $|W| = n$ cuando W es finito de n elementos.

Apoyándonos en el Teorema 14, definimos por recursión transfinita la sucesión de ordinales iniciales infinitos:

Definición 45 (Sucesión $(\omega_\alpha)_\alpha$ de ordinales iniciales)

$$\omega_0 = \omega, \quad \omega_{\alpha+1} = h(\omega_\alpha), \quad \omega_\alpha = \sup\{\omega_\beta : \beta < \alpha\} \text{ si } \alpha \text{ es límite } (\alpha \neq 0).$$

Con esta construcción, cada ω_α es un ordinal inicial infinito, y recíprocamente, todo ordinal inicial infinito coincide con algún ω_α ; la sucesión $(\omega_\alpha)_\alpha$ recoge, por tanto, exactamente y sin repeticiones todos los ordinales iniciales infinitos. En particular, todo conjunto bien ordenable es equipotente con un único ordinal inicial, y podemos identificar su cardinal con dicho representante. Esto motiva la notación estándar siguiente:

Definición 46 (Números $\aleph$) Para cada ordinal α , definimos el **número aleph** correspondiente por

$$\aleph_\alpha = \omega_\alpha,$$

en particular, $\aleph_0 = \omega_0 = \omega$.

2.7.2. Aritmética de los alephs

Trabajaremos con la aritmética cardinal introducida en la Subsección 2.5.1 (véanse Definición 36 y Proposición 9).

En particular, para $\aleph_0$ valen las identidades de la Proposición 10:

$$\aleph_0 + n = \aleph_0, \quad \aleph_0 + \aleph_0 = \aleph_0, \quad \aleph_0 \cdot \aleph_0 = \aleph_0 \quad (n \in \mathbb{N}, n > 0).$$

En el nivel transfinito, la situación es análoga: el producto de un aleph por sí mismo no aumenta su cardinalidad.

Proposición 19 (Producto de alephs) Para todo ordinal transfinito α , se cumple:

$$\aleph_\alpha \cdot \aleph_\alpha = \aleph_\alpha.$$

Demostración. Escribamos $\kappa := \aleph_\alpha$. Por definición de los números aleph, $\kappa = |\omega_\alpha|$. En [11, págs. 134–136] se construye, mediante inducción transfinita, un buen orden en $\omega_\alpha \times \omega_\alpha$ y se prueba que

$$|\omega_\alpha \times \omega_\alpha| \leq |\omega_\alpha| = \kappa.$$

Por otra parte, siempre existe una inyección $\omega_\alpha \rightarrow \omega_\alpha \times \omega_\alpha$, por ejemplo $\xi \rightarrow (\xi, 0)$, de donde se obtiene $\kappa \leq \kappa \cdot \kappa$. Aplicando el teorema de Cantor–Bernstein (Teorema 5) concluimos $\kappa \cdot \kappa = \kappa$, es decir,

$$\aleph_\alpha \cdot \aleph_\alpha = \aleph_\alpha.$$

□

Como consecuencia directa se obtiene la **regla del máximo** para cardinales infinitos: si $\alpha \leq \beta$, la suma y el producto no sobrepasan $\aleph_\beta$; en efecto,

$$\aleph_\alpha + \aleph_\beta = \aleph_\beta \quad \text{y} \quad \aleph_\alpha \cdot \aleph_\beta = \aleph_\beta.$$

Además, los finitos quedan absorbidos por cualquier aleph: para todo $n \in \mathbb{N}$ (con $n > 0$ en el caso del producto),

$$n + \aleph_\alpha = \aleph_\alpha \quad \text{y} \quad n \cdot \aleph_\alpha = \aleph_\alpha.$$

En resumen, en el ámbito de los **cardinales infinitos**, la suma y el producto quedan determinados por el mayor de los dos operandos; añadir o multiplicar por un finito no altera el tamaño una vez alcanzado el umbral infinito.

Capítulo 3

Paradojas matemáticas y su relación con la teoría de conjuntos

En el capítulo anterior hemos fijado el marco axiomático de Zermelo–Fraenkel, con y sin Axioma de Elección (ZF/ZFC), y hemos desarrollado las nociones básicas de **cardinalidad** (Sección 2.4), **ordinales** y **números aleph** (Secciones 2.6 y 2.7). Este marco proporciona un lenguaje preciso para hablar de conjuntos, funciones, relaciones y tamaños del infinito, y es el escenario en el que se formulan los resultados posteriores de este trabajo. Históricamente, sin embargo, la teoría de conjuntos nació en un contexto mucho más intuitivo, cercano a la **teoría ingenua de conjuntos**, donde se asumía que cualquier propiedad determina un conjunto. Fue precisamente en ese marco donde aparecieron una serie de **paradojas** que pusieron en tensión los fundamentos de las matemáticas y forzaron la adopción de un tratamiento axiomático más cuidadoso.

En este capítulo analizaremos tres focos paradigmáticos de inconsistencia. En primer lugar, las **paradojas de la comprensión irrestricta**, cuyo caso central es la **paradoja de Russell** (Paradoja 1), muestran que el principio de comprensión irrestricta es insostenible y conduce a distinguir entre **clases** y **conjuntos**, y a sustituirlo por el **Axioma del Esquema de separación** de Zermelo (Axioma 3). En segundo lugar, las **paradojas ligadas a la cardinalidad del infinito**, donde fenómenos como la equipotencia entre un conjunto infinito y una parte propia del mismo (Galileo, Paradoja 4), la biyección entre un intervalo y un cuadrado (Paradoja 5), la contradicción al formar un conjunto universal (Cantor, Paradoja 6), o el hotel de Hilbert (Subsección 3.2.4) ilustran el comportamiento contraintuitivo de los cardinales infinitos (Sección 2.4 y Subsección 2.5.2). Por último, las **paradojas de la buena ordenación**: la **paradoja de Burali-Forti** (Paradoja 9), que surge al intentar reunir “todos los ordinales” en un único conjunto y está estrechamente relacionada con la estructura de la clase de los ordinales y con el papel de los axiomas de **Reemplazo** y **Regularidad** (Axiomas 8 y 9 y Secciones 2.6 y 2.7), y la **paradoja de Banach-Tarski** (Paradoja 10), donde el **Axioma de Elección** (Axioma 10) permite descomposiciones geométricas con-

traintuitivas en $\mathbb{R}^3$.

Comprender cada una de estas paradojas no solo significa asomarse a los problemas que desafiaron a los matemáticos de principios del siglo XX, sino también entender por qué la teoría de conjuntos moderna necesitó una base axiomática sólida para sostenerse.

3.1. Paradojas de la comprensión irrestricta

En el Capítulo 2 hemos fijado el lenguaje y las reglas básicas de formación de conjuntos dentro de ZF/ZFC, en particular la distinción entre **conjuntos** y **clases** (Definición 3) y el papel del **Axioma del Esquema de separación** (Axioma 3) como un sustituto controlado de la comprensión ingenua. En esta sección volvemos sobre esas nociones desde una perspectiva “negativa”: estudiamos qué falla si se admite la **comprensión irrestricta** y cómo, al combinarse con definiciones autorreferenciales, aparece la **paradoja de Russell** (Paradoja 1).

3.1.1. El sistema lógico de Frege

A lo largo de su vida, Gottlob Frege, matemático y lógico alemán, formuló dos sistemas lógicos en su intento de definir los conceptos básicos de las matemáticas y derivar las leyes matemáticas a partir de las leyes de la lógica. En su libro de 1879, *Begriffsschrift* [3], desarrolló un cálculo de predicados de segundo orden que utilizó tanto para definir conceptos matemáticos como para enunciar y demostrar proposiciones de interés. Sin embargo, en su obra en dos volúmenes de 1893/1903, *Grundgesetze der Arithmetik* [4], [5], Frege amplió este sistema añadiendo como axioma lo que pensaba que era una proposición puramente lógica, la llamada **Ley Básica V**. A partir de ella intentó derivar los axiomas fundamentales de la teoría de números. Desafortunadamente, la Ley Básica V resultó no ser un principio lógico inocuo, y el sistema resultante se mostró inconsistente al estar sujeto a la paradoja de Russell.

Frege introdujo el **operador de extensión** ε , que asocia a cada concepto F un objeto εF , llamado extensión de F , entendido como “la colección de todos los objetos que caen bajo F ”.

Definición 47 (Operador de extensión) Para cada concepto F , se introduce un objeto εF (la extensión de F) tal que, en notación moderna,

$$x \in \varepsilon F \iff F(x).$$

La pieza central del sistema es la **Ley Básica V**.

Axioma 11 (Ley Básica V) Para cualesquiera conceptos F y G ,

$$\varepsilon F = \varepsilon G \iff \forall x (F(x) \iff G(x)).$$

Es decir, dos extensiones son iguales si, y sólo si, sus conceptos correspondientes coinciden en todos los objetos. La idea refleja la extensionalidad de los conjuntos (Axioma 2), pero aplicada a las extensiones de conceptos.

3.1. Paradojas de la comprensión irrestricta

Además de la Ley Básica V, en el *Grundgesetze* ([4], [5]) aparece un resultado de gran importancia conocido como el **Teorema de Frege**. Este afirma que, tomando como axioma el llamado *Principio de Hume*, es posible derivar en lógica de segundo orden los axiomas de Peano/Dedekind de la aritmética.

Axioma 12 (Principio de Hume)

$$\#F = \#G \iff \exists f : F \rightarrow G \text{ biyectiva.}$$

Es decir, el número de F -cosas es igual al número de G -cosas si y solo si existe una correspondencia biunívoca entre ambas (véase la noción de equipotencia en la Definición 31). Véase [9] para una exposición más detallada del trabajo de Frege.

A pesar de la elegancia y profundidad de este planteamiento, el sistema de Frege contenía un punto débil: la Ley Básica V, interpretada como un principio de comprensión irrestricta sobre extensiones. En 1901, Bertrand Russell descubrió que de este principio se seguía una contradicción insalvable; este hallazgo, conocido como la **paradoja de Russell**, comprometía toda la construcción de los *Grundgesetze* ([4], [5]) y motivó, en la dirección de Zermelo, la sustitución del principio de comprensión irrestricta por el **Axioma del Esquema de separación** (Axioma 3).

3.1.2. La paradoja de Russell

La contradicción descubierta por Russell puede formularse directamente en el lenguaje de conjuntos. En la teoría ingenua se asume el **principio de comprensión irrestricta**, según el cual para toda fórmula $\varphi(x)$ con x como variable libre existe el conjunto $\{x : \varphi(x)\}$.

Definición 48 (Comprensión irrestricta) Para toda fórmula $\varphi(x)$ con x como variable libre, existe el conjunto

$$\{x : \varphi(x)\},$$

cuyos elementos son exactamente los objetos que satisfacen $\varphi(x)$.

Sin embargo, si tomamos $\varphi(x)$ como “ $x \in x$ ” y definimos $R = \{x : x \notin x\}$, obtenemos el conjunto de todos los conjuntos que no son miembros de sí mismos (el llamado **conjunto de Russell**).

Paradoja 1 (Paradoja de Russell) La paradoja aparece al preguntarse si $R \in R$:

1. Si $R \in R$, entonces por definición debe cumplirse $R \notin R$.
2. Si $R \notin R$, entonces cumple la condición para pertenecer a R , con lo cual $R \in R$.

En ambos casos se llega a una contradicción. El razonamiento puede incluso formularse de manera más general, sin referencia explícita al símbolo $\in$, como muestran analogías clásicas, como la del barbero, que veremos a continuación. Esto pone de manifiesto que el problema no es exclusivo de la teoría ingenua de conjuntos, sino que refleja una dificultad lógica más profunda (véase [12] para más detalle).

Capítulo 3. Paradojas matemáticas y su relación con la teoría de conjuntos

Paradoja 2 (Paradoja del barbero) *“En cierto pueblo hay un barbero que afeita a todos, y sólo a aquellos, que no se afeitan a sí mismos. ¿El barbero se afeita a sí mismo?”*

1. Si el barbero se afeita a sí mismo, entonces es un habitante que se afeita a sí mismo, y por la definición no debería afeitarse a sí mismo.
2. Si el barbero no se afeita a sí mismo, entonces es un habitante que no se afeita a sí mismo, y por definición el barbero debería afeitarse, es decir, afeitarse a sí mismo.

De nuevo se alcanza una contradicción en ambos casos (véase [18]). El patrón lógico es el mismo que en la paradoja de Russell: un predicado definido “en términos de sí mismo” genera una situación imposible.

Paradoja 3 (Paradoja de Grelling–Nelson) *Consideremos los adjetivos “autológico” y “heterológico”. Un adjetivo es **autológico** si se describe a sí mismo. Por ejemplo, la palabra “esdrújula” es autológica (ya que es una palabra esdrújula). Un adjetivo es **heterológico** si no se describe a sí mismo. Por lo tanto, “monosílabo” es una palabra heterológica (porque tiene más de una sílaba). Aparentemente, todo adjetivo ha de ser autológico o heterológico, pues cada uno se describe a sí mismo o no. La paradoja surge al preguntar: “¿es heterológico una palabra heterológica?”*

1. Si respondemos que no, entonces “heterológico” sería autológico, lo cual significa que se aplica a sí mismo y, por tanto, debería ser heterológico.
2. Si respondemos que sí, entonces “heterológico” es heterológico, lo cual significa que no se aplica a sí mismo y, por tanto, no es heterológico.

De nuevo, en ambos casos se alcanza una contradicción (véase [13]). Tanto la **paradoja del barbero** como la de **Grelling–Nelson** ilustran, en ámbitos diferentes, la misma dificultad lógica que subyace a la paradoja de Russell; la autorreferencia no restringida conduce a contradicciones inevitables.

La paradoja de Russell pone de manifiesto que no basta formular una propiedad para garantizar la existencia del “conjunto de todos los objetos” que la satisfacen. Algunas colecciones son demasiado grandes para ser conjuntos. Para tratarlas adecuadamente conviene distinguir entre **conjuntos** y **clases**, según la terminología fijada en la Definición 3.

La respuesta fue la axiomatización de la teoría de conjuntos. La primera formulación rigurosa se debe a Ernst Zermelo (1908) (cf. Subsección 2.1.2), motivada por la necesidad de fijar principios de formación de conjuntos libres de contradicciones. En el sistema de Zermelo–Fraenkel, la paradoja de Russell se evita gracias al **Axioma de Separación** (Axioma 3), que sustituye al principio de comprensión irrestricta: ya no se postula que toda propiedad determina un conjunto, sino que sólo se permiten subconjuntos de un conjunto dado definidos por una condición.

De este modo, colecciones como la de todos los conjuntos, la de todos los ordinales o la definida por $R = \{x : x \notin x\}$, no se consideran conjuntos, sino

clases propias. La noción de clase propia permite describir estas colecciones sin poner en riesgo la consistencia del sistema, mientras que axiomas como el de Separación, Reemplazo y Regularidad (Axiomas 3, 8 y 9) garantizan que sólo se forman conjuntos a partir de otros ya existentes siguiendo reglas controladas. En este sentido, la teoría ZF (y su extensión con el Axioma de Elección, ZFC) se ha consolidado como el marco estándar para el desarrollo de las matemáticas modernas.

3.2. Paradojas de la cardinalidad del infinito

En la sección sobre conjuntos finitos, numerables y no numerables (Sección 2.4) hemos introducido la noción de **equipotencia** (Definición 31), la comparación de cardinalidades mediante inyecciones (Definición 32) y los conceptos de **conjunto finito**, **numerable** y **no numerable** (Definiciones 33 y 34), así como la no numerabilidad de $\mathbb{R}$ (Teorema 8) y la identificación del continuo con $2^{\aleph_0}$ (Teorema 11). En esta sección retomamos esos resultados desde una perspectiva más paradójica, poniendo énfasis en los choques con la intuición que motivaron su descubrimiento.

3.2.1. Galileo y los cuadrados perfectos

Galileo Galilei, físico y astrónomo italiano, señala en los *Discorsi* [6] que los números cuadrados (1, 4, 9, 16, ...) parecen ser claramente menos que los números enteros positivos (1, 2, 3, 4, ...), pues estos últimos incluyen tanto a los cuadrados como a muchos más. Sin embargo, muestra a continuación que estas dos colecciones son iguales, ya que pueden ponerse en correspondencia uno a uno (véase [16]).

En matemáticas, la manera más natural de comparar conjuntos no es contando el número de elementos de cada uno, sino estableciendo una correspondencia uno a uno entre ellos. Dicho de otro modo, dos conjuntos tienen la misma cardinalidad si existe una biyección entre sus elementos (Definición 31). Un ejemplo intuitivo consiste en pensar en un conjunto de cinco personas y un conjunto de cinco pulseras. Para comprobar que ambos conjuntos poseen la misma cardinalidad, no es necesario contar; basta con asignar a cada persona exactamente una pulsera. Si, al final del proceso, no sobra ni falta ninguna, entonces hemos construido una correspondencia biyectiva y, por tanto, los conjuntos son equipotentes.

Paradoja 4 (Paradoja de Galileo) *El conjunto de los cuadrados perfectos*

$$S := \{n^2 : n \in \mathbb{N}\}$$

es equipotente a $\mathbb{N}$, a pesar de que $S \subsetneq \mathbb{N}$.

La correspondencia se obtiene asignando a cada número natural $n \in \mathbb{N}$ su cuadrado n^2 , definiendo la aplicación $f : \mathbb{N} \rightarrow S$ con $f(n) = n^2$. La aplicación es inyectiva, pues si $f(n) = f(m)$ entonces $n^2 = m^2$ y, como $n, m \in \mathbb{N}$, se sigue $n = m$. Además, f es sobreyectiva por definición de S . Por tanto, f es biyectiva y $\mathbb{N}$ es equipotente a S .

Capítulo 3. Paradojas matemáticas y su relación con la teoría de conjuntos

A primera vista, el conjunto de los cuadrados debería tener “menos elementos”, ya que constituye una parte propia de $\mathbb{N}$. El hecho de que ambos sean equipotentes revela una característica esencial de los conjuntos infinitos: en ellos, a diferencia de los conjuntos finitos, el todo puede tener la misma cardinalidad que una de sus partes. Esta tensión entre inclusión y equipotencia es la puerta de entrada al concepto moderno de infinito numerable y al posterior desarrollo de la teoría de cardinales transfinitos de Cantor.

3.2.2. «Lo veo, pero no lo creo»

La equipotencia entre $[0, 1]$ y $[0, 1]^2$ encaja en la familia de ejemplos en los que diferentes espacios continuos comparten la **cardinalidad del continuo** (Subsección 2.5.2): en particular, $|[0, 1]| = |[0, 1]^n|$ para todo $n \geq 1$ y $|\mathbb{R}| = 2^{\aleph_0}$ (Teorema 1.1).

Richard Dedekind, matemático alemán clave en el surgimiento de la matemática conjuntista y estructural del siglo XX, recibió una carta en 1877 de Cantor en la que se le comunicaba un resultado que desafiaba la intuición geométrica: la posibilidad de establecer una biyección entre un intervalo y un cuadrado, es decir, entre conjuntos de distinta dimensión. La idea de que una línea y una superficie pudieran tener “la misma cantidad” de puntos era tan inesperada que Cantor acompañó su carta con una frase en francés que se haría célebre: *Je le vois, mais je ne le crois pas* (Lo veo, pero no lo creo) (véase [8]).

Paradoja 5 (Paradoja del intervalo–cuadrado) *Un segmento de recta y un cuadrado parecen contener “distinta cantidad” de puntos (una figura es unidimensional y la otra bidimensional); sin embargo, existe una correspondencia uno a uno entre sus puntos. En particular, el número de puntos de $[0, 1]$ coincide con el de $[0, 1]^2$, es decir, $|[0, 1]| = |[0, 1]^2|$.*

El camino seguido por Cantor para establecer la correspondencia entre el intervalo y el cuadrado comenzó con una idea ingeniosa pero problemática. Dado un par

$$(x, y) = (0, a_1 a_2 a_3 \dots, 0, b_1 b_2 b_3 \dots),$$

propuso definir

$$z = 0, a_1 b_1 a_2 b_2 a_3 b_3 \dots,$$

es decir, construir un número real intercalando las cifras decimales de x y y . De este modo, cada par de puntos del cuadrado $[0, 1] \times [0, 1]$ quedaba asociado a un único punto del intervalo $[0, 1]$. Esta construcción exige fijar un representante decimal para cada real, ya que algunos números admiten dos desarrollos: por ejemplo, $0,25000\dots = 0,24999\dots$. En lo que sigue adoptamos la convención estándar de escoger siempre la expansión que no termina en una cola infinita de nueves.

Sin embargo, Dedekind descubrió un fallo crucial: la aplicación no era sobreyectiva. Existen números del intervalo, como

$$z = 0,1201010101\dots,$$

3.2. Paradojas de la cardinalidad del infinito

que no provienen de ningún par (x, y) ya que la única posibilidad para x es $0,10000\dots$, la cual no era posible debido a la expansión decimal elegida por Cantor (la discusión histórica detallada del problema aparece en [8]).

Lejos de abandonar su idea, Cantor buscó una alternativa más sólida. Comenzó señalando que todo número irracional $x \in (0, 1)$ admite una expansión única de la forma

$$x = [0; a_1, a_2, a_3, \dots] = \frac{1}{a_1 + \frac{1}{a_2 + \frac{1}{a_3 + \ddots}}}, \quad a_i \in \mathbb{N}.$$

Si

$$x = [0; a_1, a_2, a_3, \dots], \quad y = [0; b_1, b_2, b_3, \dots],$$

entonces Cantor definía

$$z = [0; a_1, b_1, a_2, b_2, a_3, b_3, \dots],$$

es decir, una nueva fracción continua obtenida intercalando los términos de las expansiones de x e y . Como estas representaciones son únicas, se evita la ambigüedad que arruinaba la construcción con decimales. De este modo se obtiene una biyección entre los pares de irracionales $(x, y) \in [0, 1]^2$ y los irracionales de $[0, 1]$.

Para tratar los racionales, Cantor puede fijar una enumeración $\{r_k\}_{k \in \mathbb{N}}$ de $\mathbb{Q} \cap [0, 1]$, que existe porque $\mathbb{Q}$ es numerable (véase Teorema 7), y escoger una sucesión creciente de irracionales $\{\eta_k\}_{k \in \mathbb{N}} \subset (0, 1)$ con η_k convergente a 1. Considerando una permutación que intercambia $r_k \leftrightarrow \eta_k$ y fija el resto, se reduce el problema a comparar $[0, 1]$ con $[0, 1] \setminus \{\eta_k : k \in \mathbb{N}\}$. El paso técnico básico es el siguiente.

Proposición 20 (Desplazamiento en un intervalo) *Existe una biyección entre $(0, 1]$ y $[0, 1]$.*

Demostración. Definimos $f : (0, 1] \rightarrow [0, 1]$ por

$$f(1) = 0, \quad f\left(\frac{1}{n+1}\right) = \frac{1}{n} \quad (n \geq 1), \quad f(x) = x \text{ si } x \notin \left\{\frac{1}{n} : n \in \mathbb{N}\right\}.$$

Así, fuera de $\{1/n\}$ la función es la identidad y, sobre $\{1/n\}$, realiza un desplazamiento que libera el valor 0; en particular, f es biyectiva. $\square$

A partir de aquí, Cantor organiza un procedimiento iterativo que “absorbe” sucesiones numerables de puntos mediante desplazamientos de este tipo, sin afectar a los puntos ya tratados, de forma análoga a las reorganizaciones numerables descritas más adelante en el hotel de Hilbert (Subsección 3.2.4). Con ello se obtiene finalmente una biyección entre $[0, 1]$ y $[0, 1] \setminus \{\eta_k : k \in \mathbb{N}\}$ (véase [8] para la construcción histórica).

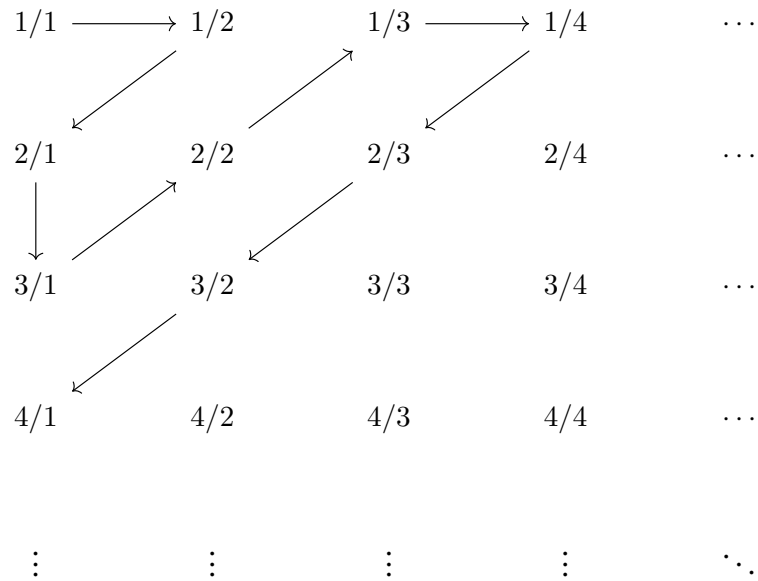
Con todo ello, Cantor obtenía una biyección rigurosa entre el intervalo $[0, 1]$ y el cuadrado $[0, 1] \times [0, 1]$. El resultado se generaliza al producto $I^n = [0, 1]^n$, de manera que

$$|[0, 1]| = |[0, 1]^n|, \quad \text{para todo } n \in \mathbb{N}.$$

Capítulo 3. Paradojas matemáticas y su relación con la teoría de conjuntos

El método diagonal. El argumento diagonal de Cantor prueba la no numerabilidad de $(0, 1)$ y, por tanto, de $\mathbb{R}$ (Teorema 8). En contraste, la construcción “en zigzag” para fracciones positivas reitera la numerabilidad de $\mathbb{Q}$ (Teorema 7).

En 1873, Cantor demuestra la numerabilidad de los números racionales, $\mathbb{Q}$. Para lograr dicha demostración empleó el método conocido como **zigzag** o recorrido diagonal de Cauchy–Cantor. Este método presenta semejanzas geométricas al posterior método diagonal. Consiste en crear una tabla que represente todas las posibles fracciones positivas, donde en cada fila va aumentando el numerador y en cada columna el denominador. Dicha demostración consiste en trazar una línea en forma de zigzag iniciando en el $1/1$.



De este modo, el recorrido en zigzag produce una enumeración de las fracciones positivas con repetición. Para obtener una enumeración de $\mathbb{Q}^+$ sin repeticiones, basta restringirse a las fracciones en forma irreducible. En efecto, sea

$$S := \{(p, q) \in \mathbb{N} \times \mathbb{N} : \text{mcd}(p, q) = 1\},$$

que es numerable por ser subconjunto de $\mathbb{N} \times \mathbb{N}$ (véase Proposición 10). En efecto, el recorrido diagonal muestra que $\mathbb{N} \times \mathbb{N}$ puede enumerarse; en términos de cardinales, $|\mathbb{N} \times \mathbb{N}| = \aleph_0$, y como $S \subseteq \mathbb{N} \times \mathbb{N}$, la inclusión $S \hookrightarrow \mathbb{N} \times \mathbb{N}$ da una inyección de S en un conjunto numerable, luego S es (a lo sumo) numerable.

La aplicación

$$S \rightarrow \mathbb{Q}^+, \quad (p, q) \mapsto \frac{p}{q},$$

es biyectiva, luego $\mathbb{Q}^+$ es numerable.

Además,

$$\mathbb{Q} = \mathbb{Q}^+ \cup \{0\} \cup (-\mathbb{Q}^+),$$

y por la aritmética de cardinales para $\aleph_0$ (Subsección 2.5.1) se tiene

$$|\mathbb{Q}| = |\mathbb{Q}^+| + 1 + |\mathbb{Q}^+|.$$

3.2. Paradojas de la cardinalidad del infinito

Como $|\mathbb{Q}^+| = \aleph_0$, esta expresión es $2\aleph_0 + 1$, y por las reglas de la Proposición 10 (añadir un número finito o tomar un número finito de copias de un numerable no cambia su cardinal), se obtiene

$$|\mathbb{Q}| = \aleph_0.$$

Por tanto, $\mathbb{Q}$ es numerable.

Una presentación clásica del **método diagonal** de Cantor aparece en 1891 al llevarse a cabo la demostración de la no numerabilidad de los números reales $\mathbb{R}$. El objetivo era comparar conjuntos infinitos discretos (como $\mathbb{N}$) con conjuntos infinitos continuos (como $(0, 1)$). En la formulación original, Cantor considera tan solo dos elementos, m y w , estableciendo el conjunto A cuyos elementos son secuencias infinitas

$$E = (x_1, x_2, x_3, \dots, x_n, \dots)$$

donde cada x_n puede ser o m o w . Supone, por absurdo, que A es numerable y escribe todos sus elementos en una lista:

$$E_1 = (a_{11}, a_{12}, a_{13}, \dots, a_{1j}, \dots)$$

$$E_2 = (a_{21}, a_{22}, a_{23}, \dots, a_{2j}, \dots)$$

$$E_3 = (a_{31}, a_{32}, a_{33}, \dots, a_{3j}, \dots)$$

$$\vdots$$

$$E_i = (a_{i1}, a_{i2}, a_{i3}, \dots, a_{ij}, \dots)$$

Es entonces cuando Cantor define el elemento $E_0 = (b_1, b_2, b_3, \dots, b_j, \dots)$, donde cada b_j se elige de modo que $b_j \neq a_{jj}$ (por ejemplo, $b_j = m$ si $a_{jj} = w$ y $b_j = w$ si $a_{jj} = m$). De esta forma, E_0 difiere de E_i al menos en la coordenada i , de modo que E_0 no coincide con ningún elemento de la lista. Esto contradice la supuesta enumeración de todos los elementos de A (véase [17] para más detalle).

Para destacar visualmente la diagonal, se puede escribir la lista como:

$$\begin{array}{lclclclcl} E_1 : & & \textcolor{red}{a_{11}} & a_{12} & a_{13} & \cdots & a_{1j} & \cdots \\ E_2 : & a_{21} & \textcolor{red}{a_{22}} & a_{23} & \cdots & & a_{2j} & \cdots \\ E_3 : & a_{31} & a_{32} & \textcolor{red}{a_{33}} & \cdots & & a_{3j} & \cdots \\ & \vdots & \vdots & \vdots & \vdots & \textcolor{red}{\ddots} & \vdots & \cdots \\ E_i : & a_{i1} & a_{i2} & a_{i3} & \cdots & & \textcolor{red}{a_{ij}} & \cdots \\ & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \textcolor{red}{\ddots} \end{array}$$

Capítulo 3. Paradojas matemáticas y su relación con la teoría de conjuntos

Por tanto, los elementos de $E_0 = (b_1, b_2, b_3, \dots, b_j, \dots)$ corresponderán a los complementarios de la diagonal de la tabla.

La formulación contemporánea del argumento de Cantor se expresa en el marco de los números reales. En lugar de trabajar con secuencias abstractas de símbolos, se consideran las expansiones decimales (o binarias) de los reales comprendidos entre 0 y 1. Supongamos, con el fin de alcanzar una contradicción, que el intervalo $(0, 1)$ es numerable. En ese caso sería posible enumerar todos sus elementos en una lista infinita $x_1, x_2, x_3, \dots$ y escribir cada uno mediante una expansión binaria

$$x_i = 0, a_{i1}a_{i2}a_{i3} \dots, \quad a_{ij} \in \{0, 1\},$$

eligiendo para cada real el representante que no termina en una cola infinita de unos (para evitar la ambigüedad $0,0111\dots = 0,1000\dots$).

$x_1 :$	0,	a_{11}	a_{12}	a_{13}	$\dots$	a_{1j}	$\dots$
$x_2 :$	0,	a_{21}	a_{22}	a_{23}	$\dots$	a_{2j}	$\dots$
$x_3 :$	0,	a_{31}	a_{32}	a_{33}	$\dots$	a_{3j}	$\dots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\dots$
$x_i :$	0,	a_{i1}	a_{i2}	a_{i3}	$\dots$	a_{jj}	$\dots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$

Usando el método diagonal se puede construir un nuevo número

$$y = 0, b_1b_2b_3 \dots,$$

donde cada cifra b_j se define distinta de a_{jj} . De esta forma, el número y difiere de x_1 en la primera cifra, de x_2 en la segunda, de x_3 en la tercera, y en general de x_j en la j -ésima. Por tanto, y no coincide con ninguno de los elementos de la lista, contradiciendo la hipótesis de que todos los números de $(0, 1)$ habían sido enumerados y, por tanto, el intervalo $(0, 1)$ no es numerable.

Dado que todo intervalo abierto es equipotente a $(0, 1)$ y que $(0, 1)$ lo es también a $\mathbb{R}$, se concluye que el conjunto de los números reales no puede ponerse en correspondencia biyectiva con los naturales, es decir, $\mathbb{R}$ no es numerable (Teoremas 8 y 11).

3.2.3. La paradoja de Cantor

La **paradoja de Cantor** constituye, junto con las de Russell (Subsección 3.1.2) y Burali-Forti (Subsección 3.3.1), una de las tres antinomias históricas fundamentales de la teoría ingenua de conjuntos. Su núcleo conceptual es la tensión entre la existencia de un **conjunto universal** y el teorema de Cantor sobre conjuntos potencia (Teorema 10), presentado en el capítulo de fundamentos.

En la teoría cantoriana ingenua se admite, vía el principio de comprensión irrestricta (Definición 48), la existencia de un conjunto U que contiene todos los conjuntos, obtenido a partir de la condición puramente lógica $x = x$:

$$U = \{x : x = x\}.$$

Este U pretende jugar el papel de “universo” de la teoría: todo objeto que pueda ser considerado conjunto pertenece a U .

Paradoja 6 (Paradoja de Cantor) *Si existe un conjunto universal U , entonces se obtiene una contradicción con el Teorema de Cantor del conjunto potencia (Teorema 10).*

Recordemos que, para todo conjunto X , la definición de conjunto potencia asocia $\mathcal{P}(X)$ como el conjunto de todos sus subconjuntos, y el resultado de Cantor garantiza estrictamente

$$|X| < |\mathcal{P}(X)|$$

en el sentido de comparación de cardinalidades introducido en la Sección 2.4. Aplicado a U , esto da

$$|U| < |\mathcal{P}(U)|.$$

Por otra parte, si U es universal, entonces todo subconjunto de U es un conjunto y, por tanto, pertenece a U ; en particular,

$$\mathcal{P}(U) \subseteq U.$$

La inclusión induce una inyección $\mathcal{P}(U) \hookrightarrow U$, y por la definición de comparación de cardinalidades (Definición 32) se sigue

$$|\mathcal{P}(U)| \leq |U|.$$

Esto contradice la desigualdad estricta $|U| < |\mathcal{P}(U)|$. Por tanto, no puede existir tal conjunto universal U .

El conflicto no está en el teorema de Cantor, que es válido para cualquier conjunto X , sino en la suposición inicial de que existe un conjunto U que contiene a todos los conjuntos, incluido su propio conjunto potencia. El análisis posterior ha mostrado que esta contradicción admite diversas reformulaciones conceptualmente equivalentes, que permiten aislar mejor su estructura lógica (véase [7]).

En todas esas variantes, el rasgo decisivo no es tanto el nombre de “conjunto universal” como la condición estructural que lo hace paradójico: que su conjunto

Capítulo 3. Paradojas matemáticas y su relación con la teoría de conjuntos

potencia siga estando contenido en él, es decir, que $\mathcal{P}(X) \subseteq X$. Cualquier colección con esta propiedad entra en conflicto con el teorema de Cantor. El caso del supuesto universo U es paradigmático, pero la misma idea está emparentada con el conjunto de Russell (Subsección 3.1.2) y con otros esquemas autorreferenciales basados en diagonalización. Así, las paradojas de Cantor y Russell aparecen como manifestaciones afines de la misma tensión entre la comprensión irrestricta y los resultados generales sobre cardinalidad desarrollados en el capítulo 2.

3.2.4. La paradoja del hotel infinito de Hilbert

El hotel de Hilbert ilustra propiedades de los conjuntos **numerables** (Subsección 2.3.1): admitir huéspedes adicionales finitos o numerables se corresponde con uniones y productos que preservan la numerabilidad, mientras que ciertos “autobuses” de huéspedes modelan conjuntos de cardinalidad estrictamente mayor, como $2^{\mathbb{N}}$, de tamaño **continuo** (Teorema 11).

En 1924, el matemático alemán David Hilbert presentó el llamado **hotel infinito**, una metáfora que con el tiempo se haría célebre. Su propósito era ilustrar, de manera accesible, las propiedades sorprendentes de los conjuntos infinitos numerables.

Paradoja 7 (Paradoja del hotel infinito) *Imaginemos un hotel con una cantidad infinita de habitaciones, numeradas $1, 2, 3, 4, \dots$, y supongamos que todas están ocupadas. A primera vista podría pensarse que no es posible admitir a nadie más; sin embargo, el comportamiento de los conjuntos infinitos permite estrategias inesperadas.*

Un huésped adicional. Si llega un nuevo huésped y todas las habitaciones están llenas, un gerente inexperto podría rechazarlo. Sin embargo, basta con pedir a cada cliente que se traslade a la habitación siguiente: el de la 1 a la 2, el de la 2 a la 3, y así sucesivamente. De este modo, la primera habitación queda libre y puede ocuparse con el nuevo cliente. Esto refleja que $\mathbb{N}$ es equipotente a $\mathbb{N} \setminus \{1\}$ y, en términos cardinales, que añadir un elemento a un conjunto numerable no cambia su cardinal:

$$\aleph_0 + 1 = \aleph_0,$$

que es el caso $n = 1$ de la Proposición 10 (1).

Un número finito de nuevos huéspedes. Si llega un autobús con n pasajeros, se desplaza a cada huésped n habitaciones hacia adelante y se alojan los recién llegados en las n primeras habitaciones libres. Esto ilustra que $\aleph_0 + n = \aleph_0$ (Proposición 10).

Un número numerable de nuevos huéspedes. Si llega un autobús con infinitos pasajeros indexados por $\mathbb{N}$, la estrategia es mover a cada huésped a la habitación cuyo número sea el doble de la que ocupaba: el de la 1 pasa a la 2, el

3.2. Paradojas de la cardinalidad del infinito

de la 2 a la 4, el de la 3 a la 6, etc. De esta manera, todas las habitaciones impares quedan disponibles, y como hay infinitas, cada pasajero del autobús puede ocupar una de ellas. Esto modela el hecho de que $\mathbb{N}$ y $2\mathbb{N}$ son equipotentes y que $2 \cdot \aleph_0 = \aleph_0$ (Proposición 10 (2)).

Infinitos autobuses con infinitos pasajeros. El hotel parece capaz de albergar a cualquiera. No obstante, la situación se vuelve aún más sorprendente cuando llega una cantidad infinita de autobuses, cada uno con infinitos pasajeros. Para resolverlo, se organiza una planilla infinita en la que se enumeran todos los pasajeros mediante pares (n, m) , donde n indica el número de autobús y m el asiento:

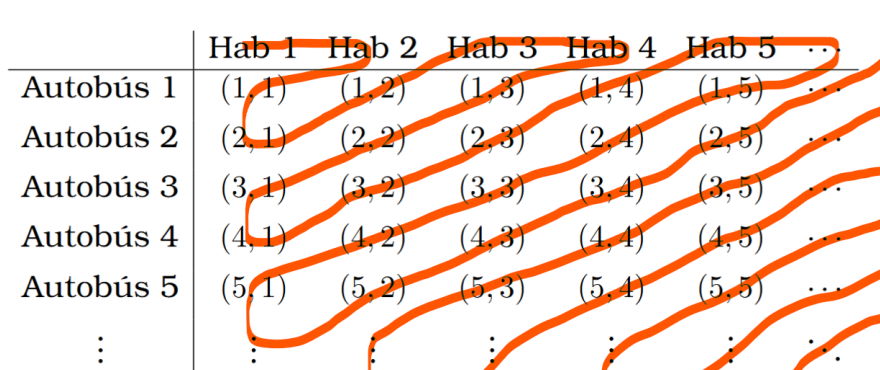
	Hab 1	Hab 2	Hab 3	Hab 4	Hab 5	...
Autobús 1	(1, 1)	(1, 2)	(1, 3)	(1, 4)	(1, 5)	...
Autobús 2	(2, 1)	(2, 2)	(2, 3)	(2, 4)	(2, 5)	...
Autobús 3	(3, 1)	(3, 2)	(3, 3)	(3, 4)	(3, 5)	...
Autobús 4	(4, 1)	(4, 2)	(4, 3)	(4, 4)	(4, 5)	...
Autobús 5	(5, 1)	(5, 2)	(5, 3)	(5, 4)	(5, 5)	...
⋮	⋮	⋮	⋮	⋮	⋮	⋮

A continuación se recorre la tabla en zigzag (Fig. 3.1), de forma análoga a la prueba de numerabilidad de $\mathbb{Q}$ (Teorema 7), obteniendo un listado lineal en el que cada par (n, m) aparece exactamente una vez y queda asociado a una habitación única. De este modo se hace explícita una biyección entre $\mathbb{N}$ y $\mathbb{N} \times \mathbb{N}$, lo que muestra que $\mathbb{N} \times \mathbb{N}$ es numerable. En términos cardinales,

$$|\mathbb{N} \times \mathbb{N}| = |\mathbb{N}|^2 = \aleph_0^2 = \aleph_0,$$

que es el caso $n = 2$ de la Proposición 10 (3).

Geométricamente, el gesto de “pellizcar” la línea por sus extremos y estirla (Fig. 3.2) representa precisamente esta idea: una planilla infinita puede reorganizarse en una única línea infinita sin perder ningún elemento.



	Hab 1	Hab 2	Hab 3	Hab 4	Hab 5	...
Autobús 1	(1, 1)	(1, 2)	(1, 3)	(1, 4)	(1, 5)	...
Autobús 2	(2, 1)	(2, 2)	(2, 3)	(2, 4)	(2, 5)	...
Autobús 3	(3, 1)	(3, 2)	(3, 3)	(3, 4)	(3, 5)	...
Autobús 4	(4, 1)	(4, 2)	(4, 3)	(4, 4)	(4, 5)	...
Autobús 5	(5, 1)	(5, 2)	(5, 3)	(5, 4)	(5, 5)	...
⋮	⋮	⋮	⋮	⋮	⋮	⋮

Figura 3.1: Representación del recorrido en zigzag del hotel infinito de Hilbert.

	Hab 1	Hab 2	(1,1)	(2,1)	(1,2)	Hab 3	...
Autobús 1							
Autobús 2							
Autobús 3							
Autobús 4							
Autobús 5							

Figura 3.2: Transformación de la planilla infinita en una línea infinita.

Un autobús “demasiado grande”. Sin embargo, existe un límite para la capacidad del hotel de Hilbert. Imaginemos un autobús especial, en el que cada pasajero está identificado no por un número, sino por un nombre único formado únicamente por las letras a y b , y que además es una secuencia infinita de dichas letras. En este caso, cada persona representa una de las infinitas sucesiones posibles de dos símbolos. El conjunto de todas esas sucesiones puede identificarse con $\{0,1\}^{\mathbb{N}}$, que tiene cardinalidad $2^{\aleph_0}$ (véase Teorema 11), estrictamente mayor que $\aleph_0$.

Al intentar organizarlos en habitaciones, se advierte que, aunque se trace una lista infinita con estos nombres, siempre es posible construir uno nuevo que no aparezca en ella, cambiando la primera letra del primer nombre, la segunda del segundo, la tercera del tercero, y así sucesivamente. Ese nuevo nombre diferirá en al menos una posición de cada uno de los de la lista, lo que muestra que no se puede alojar a todos los pasajeros de este autobús (método diagonal de Cantor).

Sin habitación	BBABA ...
Habitación 1	ABBBBBBABABBBAAAAAAAAB ...
Habitación 2	BABABBABABABABBABABAAA ...
Habitación 3	BABBBABABBBABABBABABBAB ...
Habitación 4	ABAABBBBBBAABAABABABBA ...
Habitación 5	ABABBABABAABBABBBBBBBAB ...
⋮	⋮

El hotel de Hilbert, por tanto, ilustra que, aunque existen infinitos que pueden “ampliarse” sin dejar de ser infinitos, también hay infinitos de un tamaño mayor que no pueden ser absorbidos en este proceso.

A modo de variante narrativa del mismo fenómeno, suele mencionarse el llamado diario de Tristram Shandy.

Paradoja 8 (Paradoja de Tristram Shandy) *Tristram escribe su autobiografía con tal nivel de detalle que necesita un año para relatar lo ocurrido en un solo día. Si viviera eternamente, podría, sin embargo, completar el diario: cada día de su vida*

3.3. Paradojas de la buena ordenación

puede asociarse a un año de escritura, de modo que ningún episodio quedaría sin registrar.

La aparente paradoja es puramente intuitiva. Desde el punto de vista de la teoría de conjuntos, basta observar que el conjunto de los días vividos y el conjunto de los años de escritura pueden identificarse con $\mathbb{N}$, y la correspondencia

$$n \mapsto n$$

es una biyección entre ambos (véase Definición 31). En particular, no hay contradicción: se trata del mismo fenómeno que en el hotel de Hilbert (Subsección 3.2.4), donde se reorganiza un conjunto numerable sin pérdidas, en consonancia con la aritmética de $\aleph_0$ discutida en la Proposición 10.

3.2.5. La Hipótesis del Continuo

Como resumen final de esta sección, las paradojas de cardinalidad consideradas dibujan un cuadro bastante preciso del infinito matemático. Por un lado, ejemplos como los cuadrados perfectos, el hotel de Hilbert o la correspondencia entre $[0, 1]$ y $[0, 1]^2$ muestran que muchos **conjuntos infinitos** que intuitivamente parecen de tamaños distintos resultan, sin embargo, **equipotentes**. Por otro, el método diagonal de Cantor pone de manifiesto que el conjunto de los números reales es más grande que cualquier conjunto numerable. En la escala de cardinalidades, esto se traduce en la coexistencia de al menos dos “tamaños” infinitos distintos: el infinito numerable y el infinito continuo.

En este contexto surge de manera natural la pregunta que dio origen a la **Hipótesis del Continuo (HC)**: ¿existe algún conjunto infinito cuyo tamaño esté estrictamente entre el de los números naturales y el de los números reales? Dicho de manera informal, HC afirma que no hay tal **tamaño intermedio**: todo subconjunto infinito de $\mathbb{R}$ es o bien numerable, o bien tiene la misma cardinalidad que el continuo. La formulación precisa de esta idea, en términos de la escala de los $\aleph$ y de la igualdad $2^{\aleph_0}$ con el siguiente cardinal infinito, se presentó ya en la Definición 37 y se discutió en la Subsección 2.5.2.

Históricamente, su relevancia quedó subrayada por Hilbert al situarla en primer lugar en su famosa lista de problemas (véase en detalle [10]). Hoy se sabe que la HC no puede demostrarse ni refutarse a partir de los axiomas estándar de la teoría de conjuntos (ZFC) (Subsección 2.1.2), lo que la convierte en un ejemplo paradigmático de la frontera entre lo que nuestras teorías pueden decidir y lo que permanece abierto a nuevas ampliaciones axiomáticas.

3.3. Paradojas de la buena ordenación

En el capítulo 2 se introdujeron los **ordinales** y la noción de **buen orden** como herramienta para comparar y clasificar conjuntos bien ordenados (Sección 2.6). Una idea recurrente en esa teoría es que los ordinales permiten “medir” tipos de orden, pero precisamente por su generalidad surgen distintas tensiones. En esta sección se estudian dos ejemplos paradigmáticos: la paradoja de Burali-Forti

(Paradoja 9), ligada a la imposibilidad de un “ordinal máximo”, y la paradoja de Banach–Tarski (Paradoja 10), donde la equivalencia entre elección y buena ordenación se conecta con fenómenos contraintuitivos en $\mathbb{R}^3$.

3.3.1. La paradoja de Burali–Forti

Desde un punto de vista histórico, la paradoja de Burali–Forti, publicada en 1897 por Cesare Burali–Forti, suele considerarse la primera de las paradojas lógicas modernas y constituye un punto de inflexión en el desarrollo de la teoría de conjuntos. La formulación original del autor empleaba la noción de **clase perfectamente ordenada**, que no coincidía exactamente con la noción cantoriana estándar; poco después el propio Burali–Forti reconoció la confusión y observó que el resultado paradójico se mantiene si se adopta la definición usual de buen orden en el sentido de Cantor. Por otra parte, Cantor ya había llegado hacia 1895 a la conclusión de que la colección de todos los ordinales no forma un conjunto admisible. Véase [1] para una discusión más detallada.

Paradoja 9 (Burali–Forti) *La colección de todos los ordinales no puede formar un conjunto: si se supone que existe un conjunto O cuyos elementos son todos los ordinales, se obtiene una contradicción.*

Para describir el mecanismo de la contradicción, es útil retomar el lenguaje de los ordinales introducido en la Sección 2.6. Recordemos, en particular, la noción de conjunto bien ordenado y de tipo de orden (Definición 28 y Teorema 12), así como la identificación de los ordinales con tipos de orden bien fundados (Definición 40) y la existencia de operaciones como el sucesor ordinal $\alpha + 1$ (Proposición 13). Intuitivamente, cada ordinal α representa la forma de un cierto orden bien fundado, y todo conjunto bien ordenado es isomorfo a un único ordinal.

En una primera formulación, se razona así. Puesto que Ω es un ordinal, debería pertenecer a O . Por construcción, Ω es estrictamente mayor que cualquier ordinal que aparezca en O , ya que representa el “tamaño ordinal” del conjunto completo; en particular, para todo ordinal $\alpha \in O$ se tiene $\alpha < \Omega$. Por otra parte, para cualquier ordinal γ se puede formar su sucesor $\gamma + 1$, que es estrictamente mayor que γ (Proposición 13), de manera que $\gamma < \gamma + 1$. Si aplicamos este hecho a $\gamma = \Omega$, obtenemos $\Omega < \Omega + 1$. Pero, por construcción, $\Omega + 1$ también es un ordinal y, por tanto, debería pertenecer a O , lo que implicaría $\Omega + 1 < \Omega$. Acabamos así con una desigualdad doblemente imposible:

$$\Omega + 1 > \Omega \quad \text{y} \quad \Omega + 1 < \Omega.$$

La hipótesis de que O es un conjunto conduce a una contradicción.

La formulación conceptualmente más clara se basa en la noción de **segmento inicial** (Definición 39) y en el teorema de comparación de conjuntos bien ordenados (Proposición 12). Supongamos de nuevo, por reducción al absurdo, que O es un conjunto y que su tipo de orden es Ω . Como Ω es un ordinal, también pertenece a O , y podemos considerar el segmento inicial

$$O_{<\Omega} = \{\alpha \in O : \alpha < \Omega\}.$$

Por construcción, $(O, <)$ y $(O_{<\Omega}, <)$ “enumeran” los mismos ordinales en el mismo orden, de modo que son conjuntos bien ordenados isomorfos. Pero $O_{<\Omega}$ es, por definición, un segmento inicial propio de O . Esto entra en conflicto con el Proposición 12, que afirma que para dos bien ordenados dados solo puede ocurrir una de tres posibilidades excluyentes: o bien son isomorfos, o bien uno es isomorfo a un segmento inicial propio del otro, o bien ocurre lo contrario. En nuestro caso se dan simultáneamente las dos primeras, lo cual es imposible. De nuevo, la fuente del problema es haber supuesto que la colección de todos los ordinales forma un conjunto legítimo.

Desde la perspectiva de la teoría axiomática moderna, la resolución de la paradoja es análoga a la de otras antinomias clásicas. Igual que la paradoja de Russell muestra que no toda propiedad determina un conjunto (Paradoja 1), la paradoja de Burali-Forti indica que la “colección de todos los ordinales” no puede ser un conjunto de ZF, sino una clase propia en el sentido informal de la Definición 3.

3.3.2. La paradoja de Banach-Tarski

A comienzos del siglo XX, Ernst Zermelo demostró (1904) que el **Axioma de Elección** (Axioma 10) implica que todo conjunto admite un buen orden (principio de buena ordenación). Recíprocamente, el **principio de buena ordenación** implica el Axioma de Elección; por tanto, ambos enunciados son **equivalentes**. Mucho más tarde, en 1963, Paul J. Cohen probó que en la teoría de conjuntos de Zermelo-Fraenkel sin dicho axioma no puede probarse, en general, la existencia de un buen orden para conjuntos como el de los números reales. Este contraste ilustra la enorme potencia del Axioma de Elección y anticipa fenómenos profundamente contraintuitivos cuando se combina con consideraciones geométricas en $\mathbb{R}^3$. El ejemplo paradigmático es la paradoja de Banach-Tarski, formulada por Stefan Banach y Alfred Tarski en 1924.

Paradoja 10 (Banach-Tarski) *Asumiendo el Axioma de Elección, una bola sólida en $\mathbb{R}^3$ puede descomponerse en un número finito de piezas y, mediante rotaciones y traslaciones, reensamblarse para obtener dos bolas sólidas congruentes con la original.*

Desde el punto de vista de la teoría de conjuntos desarrollada en el Capítulo 2, conviene subrayar que trabajamos dentro de $\mathbb{R}^3$, identificable con $\mathbb{R} \times \mathbb{R} \times \mathbb{R}$ (Definición 8), de manera que todos los conjuntos implicados tienen cardinalidad a lo sumo $2^{\aleph_0}$, el **cardinal del continuo** (Teorema 11). La paradoja de Banach-Tarski no introduce un “nuevo tamaño de infinito”, sino un fenómeno sorprendente sobre cómo pueden **descomponerse** y **reensamblarse** ciertos subconjuntos de $\mathbb{R}^3$ cuando se permite el uso del **Axioma de Elección**. En lo que sigue se presenta de forma esquemática el contenido de este teorema, siguiendo la exposición de [19], y se sitúa su carácter paradójico dentro del marco axiomático desarrollado en los capítulos anteriores.

El enunciado general puede describirse así. Se consideran dos subconjuntos A e B de $\mathbb{R}^3$ que estén **acotados** y tengan **interior no vacío** (respecto a la topología euclídea inducida por la métrica usual). Esto es, A está contenido en alguna bola

Capítulo 3. Paradojas matemáticas y su relación con la teoría de conjuntos

cerrada y contiene a su vez alguna bola sólida, por pequeña que sea. El resultado afirma que, bajo estas hipótesis, es posible descomponer A en un número finito de piezas $A_1, \dots, A_n$ y B en piezas $B_1, \dots, B_n$, de tal forma que cada A_j se puede trasladar y rotar rigidamente para coincidir exactamente con B_j .

Para entender el alcance de este enunciado, es útil precisar algunos términos, sin entrar en formalismos adicionales. Un **movimiento rígido** de $\mathbb{R}^3$ es una transformación que combina una rotación y una traslación, preservando distancias y ángulos; en notación vectorial puede escribirse como

$$r : \mathbb{R}^3 \rightarrow \mathbb{R}^3, \quad r(x) = \rho(x) + a,$$

donde ρ es una rotación fija y a es un vector de traslación. Dos subconjuntos A y B de $\mathbb{R}^3$ se dicen **congruentes** si existe un movimiento rígido r tal que $r(A) = B$. Más fuerte aún, A y B son **congruentes por partes** si se pueden descomponer en un número finito de piezas $A_1, \dots, A_n$ y $B_1, \dots, B_n$ de modo que, para cada j , exista un movimiento rígido r_j con $r_j(A_j) = B_j$. Nótese que detrás de esta idea está, de nuevo, la noción de biyección (Definición 31), pero ahora se exige que las biyecciones provengan de transformaciones fuertemente restringidas: movimientos rígidos de $\mathbb{R}^3$ (determinados por rotaciones y traslaciones).

El aspecto más llamativo de la paradoja aparece en el llamado resultado de “duplicación” de la bola. Se muestra que la bola unitaria cerrada B puede partitionarse en un número finito de subconjuntos $B_1, \dots, B_{40}$, y que existen movimientos rígidos $r_1, \dots, r_{40}$ tales que la familia de imágenes $\{r_k(B_k)\}$ se puede reagrupar para formar dos particiones disjuntas de B : una con 24 piezas y otra con 16. En particular, B es congruente por partes con $B \cup B$: a partir de una sola bola se obtienen dos copias congruentes entre sí y con la original. De hecho, en la demostración se controla con bastante precisión la procedencia de esas 40 piezas. En la notación técnica de [19], la partición de B no se introduce como una lista arbitraria $B_1, \dots, B_{40}$, sino como una familia doblemente refinada que surge de intersecar varias particiones construidas en pasos previos. Las piezas se denotan entonces por

$$B_{hij} \quad (1 \leq h \leq 2, 1 \leq i \leq 2, 1 \leq j \leq 10),$$

de modo que se recorren todas las combinaciones posibles y se obtienen $2 \cdot 2 \cdot 10 = 40$ subconjuntos. El índice j proviene de una partición previa de la esfera S en 10 subconjuntos T_j : dichos T_j se proyectan radialmente hacia el interior para formar las “rebanadas” de la bola (exceptuando el origen). Por su parte, los índices h e i corresponden a la elección entre dos conjuntos M_h y dos conjuntos M_i , construidos a partir de un conjunto numerable Q y del complemento $B \setminus Q$.

Con esta notación, el desglose $24 + 16$ se explica de forma directa. Las 24 **piezas** que reconstruyen la primera copia de B corresponden a las primeras seis partes de la esfera, esto es, $j = 1, \dots, 6$, combinadas con las cuatro posibilidades dadas por $(h, i) \in \{1, 2\}^2$, de manera que $6 \cdot 2 \cdot 2 = 24$. Las 16 **piezas** que reconstruyen la segunda copia se obtienen análogamente a partir de las cuatro partes restantes $j = 7, \dots, 10$, nuevamente combinadas con las mismas cuatro posibilidades para (h, i) , dando $4 \cdot 2 \cdot 2 = 16$.

3.3. Paradojas de la buena ordenación

Desde la aritmética de cardinales infinitos del Capítulo 2 (Subsección 2.7.2), tampoco hay contradicción: si denotamos por $|B|$ el cardinal de la bola, éste es un cardinal infinito, y la regla del máximo (vista como consecuencia directa de la demostración de la Proposición 19) implica que, para todo ordinal transfinito α ,

$$\aleph_\alpha + \aleph_\alpha = \aleph_\alpha.$$

En particular, es coherente que un conjunto infinito como B sea equipotente a la unión disjunta de dos copias de sí mismo. Lo verdaderamente desconcertante en la paradoja de Banach–Tarski no es esa igualdad de tamaños, sino que se obtenga mediante una descomposición en un número finito de piezas sometidas únicamente a movimientos rígidos.

La demostración del teorema se organiza en varios pasos intermedios. Primero se trabaja en la superficie de la esfera unitaria $S \subset \mathbb{R}^3$. Allí se obtiene una partición de S en un conjunto numerable P (una “parte pequeña” en el sentido de Definición 34) y tres subconjuntos S_1, S_2, S_3 con la siguiente propiedad esquemática: existen rotaciones φ, ψ_1, ψ_2 tales que

$$\varphi(S_1) = S_2 \cup S_3, \quad \psi_1(S_1) = S_2, \quad \psi_2(S_1) = S_3.$$

De manera muy intuitiva, S_1 puede “doblar” mediante rotaciones para producir conjuntos que, juntos, cubren una parte de la esfera del mismo “tamaño” que el original, y este comportamiento se traslada luego a la bola sólida extendiendo radialmente la partición de la esfera al interior de B . El resultado final es la descomposición paradójica de la bola tridimensional.

En la construcción de los conjuntos implicados aparece un paso en el que es necesario elegir, para cada órbita de puntos bajo un cierto grupo de rotaciones, un representante, sin que exista una regla explícita finita para hacerlo; esta selección global requiere precisamente el Axioma de Elección (Axioma 10), y en modelos de ZF sin elección no se puede demostrar la existencia de descomposiciones de este tipo.

Además, las piezas en las que se descompone la bola no pueden ser medibles en el sentido de Lebesgue: si todas tuvieran volumen bien definido, finito, aditivo e invariante bajo isometrías, la congruencia por partes entre una bola y dos copias de sí misma violaría inmediatamente la conservación del volumen. Además, el fenómeno es genuinamente tridimensional: en el plano $\mathbb{R}^2$ no existe un análogo por dos motivos esenciales. En primer lugar, las construcciones planas que se asemejan a una duplicación mediante rotaciones y traslaciones producen conjuntos que no satisfacen las hipótesis geométricas: el conjunto resultante es **numerable** y, sobre todo, **no acotado**. En segundo lugar, el enunciado genuino de Banach–Tarski (el Teorema H en [19]) exige que los conjuntos sean **acotados** y de **interior no vacío**. La razón estructural que subyace a esta diferencia es la complejidad del grupo de rotaciones: en $\mathbb{R}^3$ se emplea un grupo de rotaciones G generado por matrices φ y ψ con un comportamiento suficientemente “libre” como para permitir la combinatoria necesaria; en cambio, en $\mathbb{R}^2$ las rotaciones son conmutativas y no proporcionan la estructura de “palabra reducida” que hace posible una duplicación con un número finito de piezas acotadas.

Capítulo 3. Paradojas matemáticas y su relación con la teoría de conjuntos

La paradoja se sitúa así en la intersección entre los resultados sobre cardinalidades del continuo (Teorema 11), la potencia del Axioma de Elección y la geometría de $\mathbb{R}^3$, mostrando cómo ciertos principios globales de la teoría de conjuntos pueden entrar en tensión con las intuiciones habituales sobre longitud, área y volumen.

Capítulo 4

Conclusiones

4.1. Conclusiones del trabajo

El punto de partida de este trabajo ha sido la tensión entre la **teoría ingenua de conjuntos** y la formulación axiomática moderna. En este sentido, las paradojas clásicas no se han tratado como curiosidades aisladas, sino como indicadores precisos de los límites de ciertos principios de comprensión y de formación de colecciones dentro de la teoría de conjuntos.

En primer lugar, el análisis de las **paradojas de la comprensión irrestricta** y de su contexto lógico en el sistema de Frege, en particular la paradoja de Russell, ha permitido identificar con claridad el conflicto entre la comprensión irrestricta y la consistencia lógica. La introducción de la distinción entre **conjuntos** y **clases propias**, junto con el papel del **Esquema de separación**, aparece así como una respuesta estructural que delimita qué propiedades pueden emplearse para generar nuevos conjuntos y cuáles deben permanecer como descripciones de colecciones demasiado grandes.

En segundo lugar, las **paradojas de la cardinalidad del infinito** han servido para articular el comportamiento de los **cardinales infinitos** más allá de la intuición finita. Ejemplos como los cuadrados de Galileo, el hotel de Hilbert o las biyecciones entre intervalos y productos cartesianos ilustran la estabilidad del infinito numerable frente a uniones y productos, mientras que el método diagonal de Cantor y la paradoja de Cantor evidencian la existencia de cardinalidades estrictamente mayores, en particular la del **continuo**. En este marco, la **Hipótesis del Continuo** se interpreta como una cuestión sobre lo que ZFC puede decidir acerca de la estructura de la escala de cardinales entre $\aleph_0$ y $2^{\aleph_0}$.

En tercer lugar, las **paradojas de la buena ordenación** muestran hasta qué punto la teoría de ordinales refuerza la necesidad de trabajar con clases propias y aceptar fenómenos geoméricamente no intuitivos. La paradoja de Burali-Forti explica por qué la colección de todos los ordinales no puede tratarse como un conjunto sin entrar en contradicción con la propia noción de tipo de orden. La de Banach-Tarski, por su parte, ilustra el impacto del Axioma de Elección en la estructura de los subconjuntos de $\mathbb{R}^3$.

Desde el punto de vista académico, la contribución principal ha consistido en construir un marco técnico coherente para el estudio de las paradojas matemáticas relacionadas con la teoría de conjuntos. Se han recopilado y sistematizado definiciones, resultados y técnicas de la teoría axiomática de conjuntos, cardinales y ordinales, números aleph y buena ordenación, conectándolos explícitamente con los ejemplos paradójicos mediante un enfoque comparativo apoyado en fuentes clásicas y contemporáneas. Este trabajo de síntesis permite ver de manera unificada cómo los mismos principios axiomáticos intervienen en contextos lógicos, aritméticos y geométricos distintos, y sitúa las paradojas analizadas dentro de los fundamentos de la matemática contemporánea.

4.2. Análisis de impacto

Se analiza a continuación de qué manera este trabajo, centrado en las paradojas en teoría de conjuntos, se relaciona con los **Objetivos de Desarrollo Sostenible** (ODS) adoptados por las Naciones Unidas en 2015 como parte de la Agenda 2030. Dichos objetivos constituyen un marco global para orientar políticas y acciones en ámbitos tan diversos como la erradicación de la pobreza, la educación, la innovación, la igualdad, la sostenibilidad ambiental o el fortalecimiento institucional.

Aunque este Trabajo de Fin de Grado se sitúa en el terreno de la matemática teórica, ello no impide que mantenga vínculos significativos con algunos de los ODS. La formación de la matemática avanzada, la clarificación de conceptos básicos como infinito, cardinalidad o paradoja, y el desarrollo de capacidades de razonamiento abstracto resultan relevantes tanto para la educación superior como para la innovación científica y tecnológica. En el marco de la Agenda 2030, este trabajo se relaciona especialmente con los ODS 4, 9 y 17.

ODS 4: Educación de calidad. El ODS 4 busca asegurar una educación inclusiva, equitativa y de calidad y promover oportunidades de aprendizaje durante toda la vida para todos. Este trabajo contribuye a dicho objetivo principalmente a través de la formación del propio estudiante y de la creación de material que puede resultar útil en contextos docentes. Por un lado, la realización de un TFG sobre paradojas en teoría de conjuntos implica adquirir y consolidar competencias propias de la educación superior en matemáticas: razonamiento lógico, capacidad de abstracción, lectura crítica de textos especializados, redacción científica, manejo de notación formal y uso de herramientas como LaTeX. El proceso de investigación, selección de fuentes, organización del contenido y síntesis de ideas refuerza la capacidad de aprendizaje autónomo y la habilidad para comunicar resultados matemáticos de forma clara y rigurosa. Por otro lado, el texto elaborado, que revisa paradojas clásicas y las sitúa en el contexto de teorías axiomáticas modernas, puede servir como material de referencia para otros estudiantes interesados en fundamentos de la matemática, lógica o filosofía de las matemáticas. La exposición en castellano contribuye a ampliar los recursos disponibles en el ámbito hispanohablante y facilita el acceso a contenidos que, con frecuencia, se encuentran únicamente en manuales avanzados o

en lengua inglesa. De este modo, el trabajo se alinea con el ODS 4 al reforzar la calidad de la formación universitaria y al generar recursos reutilizables en asignaturas relacionadas con análisis, álgebra, lógica, teoría de conjuntos o historia y fundamentos de la matemática.

ODS 9: Industria, innovación e infraestructura. El ODS 9 propone construir infraestructuras resilientes, promover la industrialización sostenible y fomentar la innovación. Aunque el vínculo entre un estudio sobre paradojas en teoría de conjuntos y la industria pueda parecer indirecto, existen conexiones relevantes a través del papel de la matemática fundamental en la ciencia y la tecnología contemporáneas. La teoría de conjuntos, la lógica matemática y el análisis de paradojas han sido clave en el desarrollo de la informática teórica, los lenguajes de programación, los sistemas de tipos, los algoritmos de verificación de programas y, en general, en la formalización rigurosa del razonamiento. La comprensión de los límites y posibilidades de los sistemas formales resulta esencial para diseñar estructuras lógicas consistentes, sobre las que se apoyan infraestructuras digitales y tecnológicas. En este contexto, el trabajo contribuye a profundizar en la comprensión de los fundamentos que subyacen a buena parte de la matemática utilizada en ciencia e ingeniería, fomenta una cultura de rigor y de atención a la consistencia lógica. Aunque su impacto sea principalmente formativo y conceptual, se inserta en la cadena de conocimientos que sostienen la innovación científica y tecnológica, contribuyendo de manera indirecta al cumplimiento del ODS 9.

ODS 17: Alianzas para lograr los objetivos. El ODS 17 busca revitalizar la Alianza Mundial para el Desarrollo Sostenible. En el contexto de un TFG de carácter teórico, la contribución a este objetivo se manifiesta sobre todo en el plano académico e institucional. El trabajo se desarrolla en el marco de la colaboración entre estudiante y tutor, y se apoya en una red de recursos académicos (manuales, artículos, enciclopedias especializadas, materiales docentes) que son fruto de una comunidad internacional de investigadores y docentes. El uso de referencias estándar, la integración de distintas fuentes y la adaptación de contenidos para un público local ilustran cómo el conocimiento matemático es, en esencia, un esfuerzo colectivo y global. Asimismo, la posible difusión del trabajo a través de repositorios institucionales u otros canales académicos contribuye, aunque sea de manera modesta, a la creación de redes de intercambio de conocimiento coherentes con el espíritu del ODS 17.

En resumen, aunque se trata de un trabajo encuadrado en la matemática pura y en los fundamentos de la teoría de conjuntos, su impacto no es ajeno a los Objetivos de Desarrollo Sostenible. La contribución principal se articula a través de la mejora de la educación superior en matemáticas, el fortalecimiento de la base conceptual sobre la que se apoyan desarrollos científicos y tecnológicos y la participación en una red académica más amplia de producción y transmisión de conocimiento.

Bibliografía

- [1] I. M. Copi, «The Burali-Forti paradox», *Philosophy of Science*, vol. 25, n.º 4, págs. 281-286, 1958.
- [2] S. S. Epp, *Matemáticas discretas con aplicaciones*. Cengage learning, 2012.
- [3] G. Frege, *Begriffsschrift, eine der arithmetischen nachgebildete Formelsprache des reinen Denkens*. Halle a.S.: Louis Nebert, 1879, Reimpreso en *From Frege to Gödel: A Source Book in Mathematical Logic*, Harvard University Press, 1967.
- [4] G. Frege, «Grundgesetze der Arithmetik: Begriffsschriftlich abgeleitet, I», *Band, Hermann Pohle, Jena (reprinted 1962 and 1998: Georg Olms Verlagsbuchhandlung, Hildesheim)*, 1893.
- [5] G. Frege, *Grundgesetze der Arithmetik. Begriffsschriftlich abgeleitet*. Jena: Hermann Pohle, 1903, vol. II.
- [6] G. Galilei, *Discorsi e dimostrazioni matematiche*. Elsevier, 2013.
- [7] J. Garrido, «The Paradoxes of Set Theory: A Systematic Analysis», *Theoria: An International Journal for Theory, History and Foundations of Science*, págs. 35-62, 2002.
- [8] F. Q. Gouvêa, «Was cantor surprised?», *The American Mathematical Monthly*, vol. 118, n.º 3, págs. 198-209, 2011.
- [9] R. Heck y J. Burgess. «Frege's Theorem and Foundations for Arithmetic», visitado 20 de sep. de 2025. dirección: <https://plato.stanford.edu/entries/frege-theorem/>.
- [10] D. Hilbert, «Mathematical Problems», *Bulletin of the American Mathematical Society*, vol. 8, n.º 10, págs. 437-479, 1902.
- [11] K. Hrbacek y T. Jech, *Introduction to Set Theory* (Pure and Applied Mathematics), 3rd. Marcel Dekker, 1999.
- [12] A. D. Irvine y H. Deutsch. «Russell's Paradox», visitado 20 de sep. de 2025. dirección: <https://plato.stanford.edu/entries/russell-paradox/>.
- [13] R. L. Martin, «On Grelling's paradox», *The Philosophical Review*, vol. 77, n.º 3, págs. 321-331, 1968.
- [14] L. Mejía, «Peano, Lawvere, Pierce: Tres axiomatizaciones de los números naturales», Tesis doct., 2003.

- [15] J. von Neumann, «On the Introduction of Transfinite Numbers», en *From Frege to Gödel: A Source Book in Mathematical Logic, 1879-1931*, J. van Heijenoort, ed., Obra original publicada en 1923, Cambridge, Massachusetts: Harvard University Press, 1967.
- [16] M. Parker, «Philosophical method and Galileo's paradox of infinity», en *New perspectives on mathematical practices*, World Scientific, 2009, págs. 76-113.
- [17] P. Sánchez Oliva. «El método diagonal. Una aproximación a su desarrollo y aplicaciones», visitado 21 de dic. de 2025. dirección: https://www.academia.edu/88131084/El_m%C3%A9todo_diagonal_Una_aproximaci%C3%B3n_a_su_desarrollo_y_aplicaciones.
- [18] J. Siegel. «The Barber's Paradox», University of Missouri–St. Louis, visitado 12 de sep. de 2025. dirección: <https://www.umsl.edu/~siegelj/SetTheoryandTopology/TheBarber.html>.
- [19] K. Stromberg, «The Banach-Tarski paradox», *The American Mathematical Monthly*, vol. 86, n.º 3, págs. 151-161, 1979.

Anexos

Informe de originalidad (Turnitin)

De acuerdo con las recomendaciones de la Escuela, se adjunta a continuación el informe de originalidad generado mediante la herramienta Turnitin.

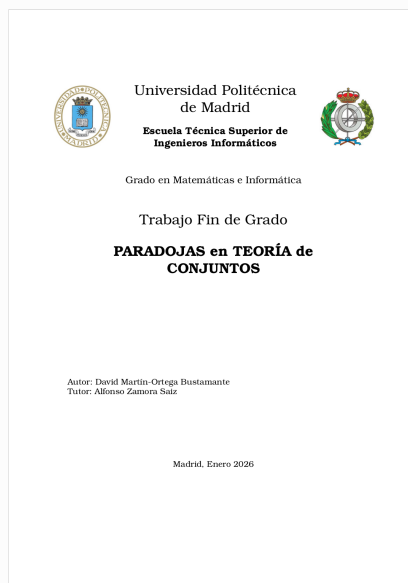


Digital Receipt

This receipt acknowledges that Turnitin received your paper. Below you will find the receipt information regarding your submission.

The first page of your submissions is displayed below.

Submission author: DAVID MARTIN-ORTEGA BUSTAMANTE
Assignment title: Turnitin Memoria Final
Submission title: Mem_Final_TFG_Paradojas_en_Teoría_de_Conjuntos.pdf
File name: 22653_DAVID_MARTIN-ORTEGA_BUSTAMANTE_Mem_Final_TF...
File size: 725.26K
Page count: 66
Word count: 23,809
Character count: 116,484
Submission date: 15-Jan-2026 01:45PM (UTC+0100)
Submission ID: 2857244975



Este documento esta firmado por



Firmante	CN=tfgm.fi.upm.es, OU=CCFI, O=ETS Ingenieros Informaticos - UPM, C=ES
Fecha/Hora	Thu Jan 15 15:23:34 CET 2026
Emisor del Certificado	EMAILADDRESS=camanager@etsiinf.upm.es, CN=CA ETS Ingenieros Informaticos, O=ETS Ingenieros Informaticos - UPM, C=ES
Numero de Serie	561
Metodo	urn:adobe.com:Adobe.PPKLite:adbe.pkcs7.sha1 (Adobe Signature)