



LA COINDUCCIÓN MATEMÁTICA EN LA CONSTRUCCIÓN DE LOS NÚMEROS REALES

JAMES ADRIÁN QUINTERO PÉREZ
AIRON STIVEN CASTIBLANCO MONTAÑEZ

UNIVERSIDAD DEL VALLE
INSTITUTO DE EDUCACIÓN Y PEDAGOGÍA
LICENCIATURA EN EDUCACIÓN BÁSICA CON ÉNFASIS EN
MATEMÁTICAS
SANTIAGO DE CALI

2014



LA COINDUCCIÓN MATEMÁTICA EN LA CONSTRUCCIÓN DE LOS NÚMEROS REALES

JAMES ADRIÁN QUINTERO PÉREZ
AIRON STIVEN CASTIBLANCO MONTAÑEZ

**Trabajo de Grado para optar por el título de
LICENCIADO EN EDUCACIÓN BÁSICA CON ÉNFASIS EN
MATEMÁTICAS**

Dirigen:

MARIBEL PATRICIA ANACONA
INSTITUTO DE EDUCACIÓN Y PEDAGOGÍA

GUILLERMO ORTÍZ
DEPARTAMENTO DE MATEMÁTICAS

UNIVERSIDAD DEL VALLE
INSTITUTO DE EDUCACIÓN Y PEDAGOGÍA
LICENCIATURA EN EDUCACIÓN BÁSICA CON ÉNFASIS EN
MATEMÁTICAS
SANTIAGO DE CALI

2014

Santiago de Cali, 12 de Septiembre de 2014

Nota de aceptación

Director 1

Director 2

Evaluador 1.

Evaluador 2.

AGRADECIMIENTOS

Expresamos nuestro agradecimiento a Dios por brindarnos la salud, el espacio y el tiempo para desarrollar nuestra tesis, de igual forma agradecemos a nuestros familiares que nos apoyaron en el transcurso de nuestra carrera.

Agradecemos a los profesores Guillermo Ortíz y Maribel Anacona por brindarnos las pautas y hacernos madurar intelectualmente, además por ser una luz que enmarca un sendero para nuestra futura formación.

Agradecemos a nuestros compañeros y amigos que creyeron en nosotros y nos brindaron su ayuda. Por último y no menos importante agradecemos a los profesores Octavio Pabón y Cesar Ojeda (Q.D.P.), por formarnos como seres humanos y darnos herramientas para defendernos en el mundo laboral.

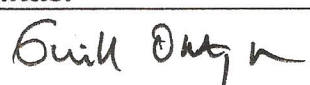
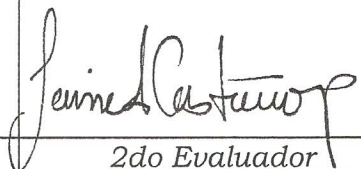


Acta de Evaluación de Trabajo de Grado

Tenga en cuenta: 1. Marque con una **X** la opción escogida.
2. diligencie el formato con una letra legible.

Título del Trabajo:	La coinducción matemática en la construcción de los números reales					
Se trata de:	Proyecto ☐		Informe Final ☒			
Directores:	Guillermo Ortiz Rico Maribel Anacona					
1er Evaluador:	Luz Victoria De La Pava					
2do Evaluador:	Jaime Andrés Castaño					
Fecha y Hora:	Año:		Mes:		Día:	Hora:
Estudiantes						
Nombres y Apellidos completos			Código		Programa Académico	
JAMES ADRIAN QUINTERO PEREZ			200935645		3469	
AIRON STIVEN CASTIBLANCO MONTAÑEZ			200942858		3469	

Evaluación					
Aprobado	☒	Meritorio	☐	Laureado	☒
Aprobado con recomendaciones	☐	No Aprobado	☐	Incompleto	☐
En el caso de ser Aprobado con recomendaciones (diligenciar la página siguiente), éstas deben presentarse en un plazo de _____ (máximo un mes) ante:					
Director del Trabajo		☐	1er Evaluador		☐
2do Evaluador		☐			
En el caso que el Informe Final se considere Incompleto , se da un plazo de máximo de _____ semestre(s) para realizar una nueva reunión de evaluación el:					
Año:		Mes:		Día:	Hora:
En el caso que no se pueda emitir una evaluación por falta de conciliación de argumentos entre Director, Evaluadores y Estudiantes; expresar la razón del desacuerdo y las alternativas de solución que proponen (diligenciar la página siguiente).					

Firmas:		
 		
Director del Trabajo de Grado	1er Evaluador	2do Evaluador



PARTE 1. Términos de la licencia general para publicación digital de obras en el repositorio institucional de Acuerdo a la Política de Propiedad Intelectual de la Universidad del Valle

Actuando en nombre propio los AUTORES o TITULARES del derecho de autor confieren a la UNIVERSIDAD DEL VALLE una Licencia no exclusiva, limitada y gratuita sobre la obra que se integra en el Repositorio Institucional, que se ajusta a las siguientes características:

a) Estará vigente a partir de la fecha en que se incluye en el Repositorio, por un plazo de cinco (5) años, que serán prorrogables indefinidamente por el tiempo que dure el derecho patrimonial del AUTOR o AUTORES. El AUTOR o AUTORES podrán dar por terminada la licencia solicitando por escrito a la UNIVERSIDAD DEL VALLE con una antelación de dos (2) meses antes de la correspondiente prórroga.

b) El AUTOR o AUTORES autorizan a la UNIVERSIDAD DEL VALLE para que en los términos establecidos en el Acuerdo 023 de 2003 emanado del Consejo Superior de la Universidad del Valle, la Ley 23 de 1982, Ley 44 de 1993, Decisión Andina 351 de 1993 y demás normas generales sobre la materia, publique la obra en el formato que el Repositorio lo requiera (impreso, digital, electrónico, óptico, usos en red o cualquier otro conocido o por conocer) y conocen que dado que se publica en Internet por este hecho circula con un alcance mundial.

c) El AUTOR o AUTORES aceptan que la autorización se hace a título gratuito, por lo tanto renuncian a recibir emolumento alguno por la publicación, distribución, comunicación pública y cualquier otro uso que se haga en los términos de la presente Licencia y de la *Licencia Creative Commons* con que se publica.

d) El AUTOR o AUTORES manifiestan que se trata de una obra original y la realizó o realizaron sin violar o usurpar derechos de autor de terceros, obra sobre la que tiene (n) los derechos que autoriza (n) y que es él o ellos quienes asumen total responsabilidad por el contenido de su obra ante la UNIVERSIDAD DEL VALLE y ante terceros. En todo caso la UNIVERSIDAD DEL VALLE se compromete a indicar siempre la autoría incluyendo el nombre del AUTOR o AUTORES y la fecha de publicación. Para todos los efectos la UNIVERSIDAD DEL VALLE actúa como un tercero de buena fé.

e) El AUTOR o AUTORES autorizan a la UNIVERSIDAD DEL VALLE para incluir la obra en los índices y buscadores que estimen necesarios para promover su difusión. El AUTOR o AUTORES aceptan que la UNIVERSIDAD DEL VALLE pueda convertir el documento a cualquier medio o formato para propósitos de preservación digital.

SI EL DOCUMENTO SE BASA EN UN TRABAJO QUE HA SIDO PATROCINADO O APOYADO POR UNA AGENCIA O UNA ORGANIZACIÓN, CON EXCEPCIÓN DE LA UNIVERSIDAD DEL VALLE, LOS AUTORES GARANTIZAN QUE SE HA CUMPLIDO CON LOS DERECHOS Y OBLIGACIONES REQUERIDOS POR EL RESPECTIVO CONTRATO O ACUERDO.

PARTE 2. Autorización para publicar y permitir la consulta y uso de obras en el Repositorio Institucional.

Con base en este documento, Usted autoriza la publicación electrónica, consulta y uso de su obra por la UNIVERSIDAD DEL VALLE y sus usuarios de la siguiente manera;

a. Usted otorga una (1) licencia especial para publicación de obras en el repositorio institucional de la UNIVERSIDAD DEL VALLE (Parte 1) que forma parte integral del presente documento y de la que ha recibido una (1) copia.

Si autorizo ☒ No autorizo ☐

b. Usted autoriza para que la obra sea puesta a disposición del público en los términos autorizados por Usted en los literales a), y b), con la **Licencia Creative Commons Reconocimiento - No comercial - Sin obras derivadas 2.5 Colombia** cuyo texto completo se puede consultar en <http://creativecommons.org/licenses/by-nc-nd/2.5/co/> y que admite conocer.

Si autorizo ☒ No autorizo ☐

Si Usted no autoriza para que la obra sea licenciada en los términos del literal b) y opta por una opción legal diferente descríbalala¹:

En constancia de lo anterior,

Título de la obra: La Coinducción Matemática en la Construcción de los Números Reales.

Autores:

Nombre: Airon Stiven Castiblanco Montañez.

Firma: Airon Castiblanco
C.C. 1144040518

Nombre: James Adrián Quintero Pérez.

Firma: James Quintero
C.C. 1143936024

Nombre:

Firma:
C.C.

Fecha: 14-10-2014

¹ Los detalles serán expuestos de ser necesario en documento adjunto

RESUMEN

El presente trabajo de grado tiene como objeto de interés ilustrar la manera en que el principio de coinducción posibilita una vía alternativa para el estudio ó construcción de los números reales, lo cual pone de manifiesto que las matemáticas como ciencia no son siempre es inductivas. Se estudia a $\mathbb{R}$ mediante la dualidad que existe con $\mathbb{N}$ por medio de la teoría de categorías, la cual permite observar la estructura de estos objetos y definir a los números naturales como un objeto inicial, y por lo tanto cumple con los principios de inducción y de recursión. Dualmente, los números reales se definen como un objeto final, en consecuencia se establecen los principios de coinducción y de corecursión.

Palabras Clave: Números reales, números naturales, principio de inducción, principio de recursión, principio de coinducción, principio de corecursión.

Índice general

Introducción	1
1. PRELIMINARES	5
1.1. Inducción y Recursión	5
1.2. Axiomática de Conjuntos	6
1.3. El Sucesor	14
1.4. Los Números Naturales	14
1.5. Principio de Inducción	17
1.5.1. Orden en los Naturales	19
1.5.2. Principio de Recursión	21
1.5.3. Operaciones Con Números Naturales	23
2. Inducción y Coinducción	26
2.1. Teoría de puntos fijos	26
2.2. Retículos	30
2.3. Principio De Coinducción	37
3. Dualidad Entre $\mathbb{N}$ y $\mathbb{R}$ en Categorías	41
3.1. Categorías	42
3.2. Álgebras Tipo 0 - 1	48
3.3. $\mathbb{N}$ como Objeto Inicial en Teoría de Categorías	49
3.3.1. Principio De Inducción	53
3.3.2. El principio de recursión.	55
3.4. Coálgebra y Coinducción	56
3.4.1. $\mathbb{R}$ como Objeto Final y Coálgebra Final	62
3.5. La Ecuación de Pavlović	69
3.6. La Coinducción en el Cálculo	71
3.6.1. La Serie de Taylor desde La Coinducción	76

IV

ÍNDICE GENERAL

Conclusión

83

Introducción

El presente trabajo de grado se inscribe en la línea de formación en Historia y Epistemología de las Matemáticas; del Programa de Licenciatura en Educación Básica con énfasis en Matemáticas del Instituto de Educación y Pedagogía de la Universidad del Valle; y está dirigido a estudiantes de licenciatura en matemáticas de último año y a docentes de matemáticas que ejerzan y deseen estudiar los números reales desde otra perspectiva.

En lo que se refiere al estudio de los números reales, es necesario tener en cuenta el proceso de construcción a través de la historia; por consiguiente, en el siglo XIX buena parte de las matemáticas se dedican a la fundamentación del análisis, y el ideal para lograr este objetivo fue la construcción de un conjunto numérico continuo que prescindiera de la noción de magnitud. Como resultado se obtuvieron diversas construcciones de los números reales, cada una con su respectiva noción de continuidad y completitud.

Tradicionalmente se estudian dos construcciones de los números reales surgidas en el siglo XIX: a través de cortaduras por Richard Dedekind y a través de sucesiones de Cauchy por Georg Cantor, de las cuales se hablará un poco a continuación. Cabe aclarar que en este trabajo no se hará un estudio histórico de estas dos construcciones, si no que se tomarán algunos resultados que ayudarán a evidenciar la dualidad que existe entre $\mathbb{N}$ y $\mathbb{R}$.

Dedekind tiene como objetivo construir un conjunto numérico continuo ó completo y para esto parte de un conjunto base, como son los racionales, en este sentido los racionales no constituyen una construcción si no una creación a partir de los enteros. Dedekind¹ para cumplir su objetivo identifica tres momentos: en el primero toma a $\mathbb{Q}$ como un cuerpo numérico ordenado. En

¹La construcción de los números reales por cortaduras de Dedekind se encuentra en [6]

el segundo momento observa que los números racionales generan cortaduras, siendo estas denominadas como cortaduras de tipo uno. Por otro lado, Dedekind observa que existen las cortaduras que son generadas por números no racionales, a estas las denomina cortadura de tipo dos. Finalmente Dedekind une los dos tipos de cortaduras para obtener un conjunto numérico continuo llamado $\mathbb{R}$.

Cantor en 1872 fue el primero en publicar una construcción² de los reales con sucesiones de Cauchy, para esto inició con los números racionales con su estructura de cuerpo numérico ordenado. Luego pensó en cada racional no como un número si no como una sucesión de Cauchy; es decir, los números reales se representan como sucesiones de Cauchy. Posteriormente, hizo paquetes o colecciones de sucesiones que tuvieran el mismo límite formando así clases de *equivalencia*. A estas clases de equivalencia de sucesiones de Cauchy se les define una estructura básica como la suma, el producto y el orden, para formar un cuerpo numérico ordenado el cual contiene una copia exacta de $\mathbb{Q}$ y tiene exactamente las mismas propiedades algebraicas y de orden de $\mathbb{Q}$. Además cumple con la característica de que cada sucesión de Cauchy coincide con las sucesiones convergentes, entonces se obtiene un nuevo cuerpo numérico ordenado que contiene los límites que le faltaban a las sucesiones de Cauchy en $\mathbb{Q}$, denominado $\mathbb{R}$.

Por otro lado, al inicio del siglo XX, la naturaleza de las matemáticas es esencialmente axiomática y formalista, las cuales se basaron en las propuestas estructurales que surgieron a finales del siglo XIX. La estructuralización de las matemáticas, se puede dividir en tres momentos:

- El primero surgió con la matemática moderna a mediados del siglo XIX. Hilbert en 1900 en su libro de *Fundamentos de las Matemáticas* propone el método axiomático, el cual consiste en dar existencia a los objetos matemáticos a través de constructos teóricos previamente definidos, denominados axiomas. Según Zalamea en [24] el método axiomático, pertenece a la matemática moderna y se fundamenta en axiomas y definiciones que se ligan. Por ejemplo, la teoría de conjuntos que se fundamenta con la lógica matemática, esto permite definir un objeto matemático como los números reales.

²Esta construcción se encuentra en el artículo de Cantor *Sobre la extensión de un teorema de la teoría de las series trigonométricas* en 1872.

- Un segundo momento transcurre con los miembros del grupo Bourbaki en los años 1940-1960, cuando presentan su programa unificador de las matemáticas introduciendo las estructuras madres (algebraica, de orden y topológica) en el interior de las axiomáticas básicas. Esta propuesta metodológica no solo ocupa un lugar de privilegio en la historia reciente de las matemáticas sino que es reconocida por su enorme influencia en la formación matemática universitaria.
- Un tercer momento aparece con el estructuralismo categórico de Eilenberg, Mac Lane y Grothendieck a partir de los años 1960. También llamado las matemáticas contemporáneas, que se fundamentan en el método sintético, el cual es un proceso que relaciona hechos aparentemente aislados y se formula una teoría que unifica los diversos elementos; es decir, consiste en la reunión racional de varios elementos dispersos en una nueva totalidad.

Este trabajo hace uso de conceptos desarrollados en las matemáticas contemporáneas, sin excluir los fundamentos históricos de los otros momentos, y propone otra forma de ver a los números reales, por medio del método conocido como el principio de coinducción, el cual permite estudiar la estructura de los objetos.

El objetivo principal de este trabajo es evidenciar la importancia del principio de coinducción, que permite comprender desde otra perspectiva la construcción de los números reales. Para esto se identifica la dualidad que existe entre la inducción y la coinducción. El primer principio es el método de demostración de conjuntos inductivos como el conjunto de los números naturales; dualmente, el principio de coinducción es el método de demostración de los conjuntos coinductivos como el conjunto de los números reales. Lo anterior permite decantar una dualidad entre $\mathbb{N}$ y $\mathbb{R}$, fundamental para la comprensión de $\mathbb{R}$.

Los números naturales se estudian por medio de la inducción y la recursión, teniendo como fundamento los axiomas de Zermelo-Fraenkel, con el axioma de elección (ZFC). Esto con el fin de realizar un estudio epistemológico de la construcción de dicho conjunto e identificarlo como un objeto y un álgebra inicial. De forma dual se estudia a los números reales como una coálgebra final y como un objeto final que tiene inmerso los principios de

coinducción y corecursión. Estos principios constituyen la forma de razonar sobre las coálgebras y son duales a los principios de inducción y recursión. Estos conceptos con los que se estudia a $\mathbb{R}$ surgen de una teoría axiomática consistente que no incluye el axioma de fundamentación, que es la de Peter Aczel [1]. En consecuencia, con el fin de comprender mejor el concepto de coinducción se hace un estudio desde la teoría de puntos fijos, retículos y teoría de categorías.

Esta última teoría es un lenguaje usado en la matemática desde mediados del siglo XX, y es la generalización de la teoría de retículos³, puesto que, las categorías son la generalidad de los conjuntos parcialmente ordenados, de igual forma sucede con los conceptos de: funtor, álgebras y coálgebras que son la generalización de función monótona, puntos prefijos y puntos postfijos, respectivamente.

Este trabajo de grado se ha estructurado en tres capítulos:

En el primer capítulo, se presentan los números naturales desde la axiomática de ZFC. Además se explican los principios de inducción y de recursión con el fin de identificar a $\mathbb{N}$ como el menor conjunto inductivo y como un objeto inicial.

En el segundo capítulo, se presentan los principios de coinducción y corecursión como una dualidad a los principios de inducción y recursión, por medio de la teoría de puntos fijos y retículos, con el fin de definir los conjuntos inductivos y coinductivos.

En el tercer y último capítulo, se presenta la dualidad entre $\mathbb{N}$ y $\mathbb{R}$ mediante la teoría de categorías, con el fin de identificar a $\mathbb{N}$ como un álgebra inicial y a $\mathbb{R}$ como una coálgebra y objeto final. Además, esta dualidad permite exhibir evidencias que ilustran el carácter coinductivo del análisis en contraste con el carácter inductivo de la aritmética.

³El concepto de retículo permite distinguir las propiedades de una clase de álgebra con otra. La teoría de retículos surgió desde el álgebra de Boole's y el estudio de divisibilidad de Dedekind, y comenzó a cobrar impulso en la década de 1930 siendo en gran medida promovida por el libro de *Lattice Theory* en la década de 1940 [3].

Capítulo 1

PRELIMINARES

En este capítulo se presentan los principios de inducción y recursión, que definen los números naturales por ser un objeto inicial de acuerdo con [17]. Con el fin de presentar la forma clásica en la cual se construyen los números naturales en la teoría de conjuntos, se presentan los axiomas de ZFC, siendo necesarios para comprender la construcción de los números naturales, y así evidenciar las propiedades que permiten caracterizar a los números naturales como objetos iniciales.

Al identificar los principios de inducción y recursión de los números naturales, se puede establecer la dualidad existente entre $\mathbb{N}$ y $\mathbb{R}$, que se estudiará en los capítulos dos y tres.

1.1. Inducción y Recursión

La inducción es un proceso que permite demostrar propiedades sobre objetos que están definidos recursivamente, un ejemplo de ello son los números naturales que se definen recursivamente así:

- I. $0 \in \mathbb{N}$
- II. Si $n \in \mathbb{N}$ entonces el *sucesor de n* pertenece a $\mathbb{N}$ y se denota como $suc(n)$
- III. $\mathbb{N}$ es el menor conjunto que cumple con las propiedades anteriores.

El principio de inducción es un método de demostración, y la recursión es un método que garantiza la buena definición de las operaciones entre objetos inductivos. Cuando se realiza una demostración por el método de inducción, implícitamente se está utilizando el principio de recursión, pues éste garantiza que las operaciones básicas de la suma y el producto estén bien definidas.

Por ejemplo, los números naturales se forman recursivamente puesto que, pueden ser formados a partir del número 0 y dado cualquier n , se puede obtener el sucesor. Al realizar las operaciones suma o producto con dos elementos de los números naturales el resultado es otro número natural.

Observación 1.1 *La diferencia entre el principio de inducción y recursión es que el primero es un método de demostración sobre propiedades del conjunto de los números naturales y el segundo permite establecer la buena definición de las operaciones sobre $\mathbb{N}$.*

Hasta el momento se ha definido de una manera muy intuitiva y general los principios de inducción y recursión; sin embargo falta aún profundizar más sobre estos, debido a que son fundamentales para entender la dualidad que existe entre $\mathbb{N}$ y $\mathbb{R}$. Para definir formalmente los dos principios es necesario definir a los números naturales y para esto escogeremos partir de una axiomática.

Para comprender las propiedades de los números naturales es necesario reconstruir a este objeto matemático y para ello se empleará el sistema axiomático de (ZFC). Este método es una forma básica de construcción de los números naturales en la teoría de conjuntos, que permite identificar a este conjunto numérico como el más pequeño de los conjuntos inductivos; es decir, como un objeto minimal.

1.2. Axiomática de Conjuntos

El sistema ZFC está formado por los siete axiomas propuestos por Zermelo en 1908, por el axioma de reemplazo introducido por Fraenkel y Skolem en 1922 y por el axioma de fundamentación incluido posteriormente por Zermelo en su “sistema ZF extendido” [26]. Los axiomas de Zermelo son el de extensionalidad, de los conjuntos elementales, de separación, conjunto potencia,

de la unión, de elección y del infinito [25]. A continuación presentamos estos axiomas tal como se exhiben generalmente en los textos modernos de teoría de conjuntos. Su incorporación en el documento responde a la necesidad de que este trabajo sea auto contenido. Sin embargo si el lector lo considera pertinente, puede avanzar a la lectura de la sección siguiente.

Los axiomas¹ de ZFC son:

[1]Axioma de Extensión

Un conjunto A es igual a un conjunto B, si y sólo si, tienen los mismos elementos. Es decir cuando todo elemento de A es un elemento de B y todo elemento de B es un elemento de A.

$A = B$, si y sólo si, $(x \in A \text{ implica } x \in B)$ y $(x \in B \text{ implica } x \in A)$.

Dos conjuntos son iguales si tienen los mismos elementos sin importar si se repiten varias veces. Por ejemplo sean los conjuntos A y B:

$A = \{a, b, c\}$, y $B = \{a, a, b, c, c\}$

Con lo anterior los conjuntos A y B son iguales.

De acuerdo a este axioma se tiene los siguientes resultados:

I. $A = A$

II. $A = B \text{ implica } B = A$

III. $A = B \text{ y } B = C \text{ implica } A = C$

Definición 1.1 *B es un subconjunto de A, si y sólo si: $\forall x (x \in B \text{ implica } x \in A)$. Que usualmente se resume por $B \subseteq A \leftrightarrow (x \in B \rightarrow x \in A)$, donde el cuantificador sobre la variable x se sobreentiende.*

¹Estos axiomas son tomados de [12] y como ayuda para una mejor comprensión se recomienda [16]

[2] Axioma de Especificación ó de Separación

Si A es un conjunto, existe otro conjunto contenido en A , que satisface una propiedad determinada. En otras palabras si se parte de un conjunto A determinado, se obtiene otro conjunto B , el cual es una parte de A , y cuyos elementos cumplen una propiedad ² $S(x)$.

Una definición más precisa de este axioma es: Dado un conjunto A y una propiedad $S(x)$, existe un conjunto B cuyos elementos son aquellos elementos de A que satisfacen la propiedad $S(x)$. Por el axioma de extensión, B es único y se denota por $\{x \in A : S(x)\}$.

Este axioma permite obtener nuevos conjuntos de un conjunto dado, a partir de una propiedad cualquiera, y tiene como fin limitar el tamaño de los conjuntos que se pueden formar lícitamente.

Hasta aquí no se ha introducido la existencia de un conjunto alguno, es por ello que el siguiente axioma lo permite y en algunos casos se opta por empezar con este.

[3] Axioma del Conjunto Vacío

Existe un conjunto el cual no contiene elementos. Es decir, $(\exists A)(\forall x) (x \notin A)$

Proposición 1.1 *Existe un único conjunto vacío.*

Demostración

Sea A un conjunto vacío, el cual existe por el axioma de existencia. Supongamos que B es otro conjunto vacío, entonces la proposición $x \in A \leftrightarrow x \in B$ es verdadera, ya que las implicaciones $x \in A \rightarrow x \in B$ y $x \in B \rightarrow x \in A$, son verdaderas puesto que sus antecedentes son falsos ■

Proposición 1.2 *Para todo conjunto A se tiene que: $\emptyset \subseteq A$*

²Formalmente una propiedad es una fórmula de primer orden con ocurrencia de variables libres (no cuantificadas), por ejemplo si X es una variable libre en la fórmula se dice una propiedad sobre X .

Demostración

$$x \in \emptyset \rightarrow x \in A$$

Todo elemento de $\emptyset$ es un elemento de A , pero como $\emptyset$ carece de elementos, entonces, $x \in \emptyset$ es una proposición falsa y por lo tanto la implicación es verdadera. En consecuencia la condición se cumple de manera inmediata. ■

[4] Axioma de Pares

Dados dos conjuntos cualesquiera, existe un conjunto que los contiene a ambos. Este axioma permite ampliar los conjuntos a partir de dos ya existentes.

$\forall A \forall B \exists C (x \in C \leftrightarrow x = A \vee x = B)$. Por el axioma de extensión C es único y se denota $\{A, B\}$

Ejemplo

Sea $A = \emptyset$ y $B = \emptyset$ entonces $C = \{\emptyset, \emptyset\} = \{\emptyset\}$ (por el axioma de extensión).

Este axioma permite obtener $D = \{\emptyset, \{\emptyset\}\}$

Observación 1.2 $\emptyset \neq \{\emptyset\}$. Dado que $\emptyset$ es un conjunto sin elementos y $\{\emptyset\}$ es un conjunto con el elemento vacío.

[5] Axioma de la Unión

Dado dos conjuntos A y B , se forma la colección cuyos elementos son todos los elementos de A , junto con los elementos de B , que se designa por $A \cup B$.

De esta forma $x \in A \cup B \leftrightarrow x \in A \vee x \in B$

[6] Axioma de las Uniones

La unión arbitraria de conjuntos es un conjunto.

Sea $\mathcal{C}$ una colección de conjuntos, el axioma anterior garantiza la existencia de la colección $\bigcup \mathcal{C}$. Es decir $x \in \bigcup \mathcal{C} \leftrightarrow \exists X \text{ tal que } X \in \mathcal{C} \wedge x \in X$

Observación 1.3 $\bigcup \emptyset = \emptyset$

Definición 1.2 *La intersección de dos conjuntos A y B , se define como el conjunto formado por los elementos comunes a A y B .*

$$x \in A \cap B \leftrightarrow x \in A \wedge x \in B.$$

Intersecciones Arbitrarias

Sea $\mathcal{C}$ una colección no vacía de conjuntos, se define la intersección generalizada así:

Definición 1.3 *la intersección generalizada de $\mathcal{C}$, una colección no vacía de conjuntos con $C \in \mathcal{C}$, está dada por*

$$\bigcap \mathcal{C} = \{x \in C : x \in X, \text{ para todo } X \in \mathcal{C}\}$$

Observación 1.4 $\bigcap \emptyset$ no se puede definir porque el conjunto vacío carece de elementos, y una $\bigcap \emptyset$ correspondería a la colección de todos los conjuntos.

[7] Axioma De Las Potencias

Dado un conjunto A entonces *partes de A* es conjunto. Y esta dado por la colección de todos los subconjuntos de A . Se denota por $\wp(A)$.

$$\forall A \exists B \forall x (x \in B \leftrightarrow x \subseteq A)$$

Ejemplo Si el conjunto $A = \{a, b\}$, entonces:

$$\wp(A) = \{\emptyset, \{a\}, \{b\}, \{a, b\}\}$$

[8] Axioma Del Infinito

En 1870 Georg Cantor da el sentido moderno del tamaño de un conjunto a través de las biyecciones y proporciona el primer estudio sistemático de los conjuntos, los números ordinales y cardinales.

Notación: la cardinalidad de un conjunto X se denota por $\text{Card}(X)$

$$\wp(\emptyset) = \emptyset \quad \text{Card}(A) = 0 \rightarrow \text{Card}(\wp(A)) = 1$$

$$\wp(a) = \{\emptyset, a\} \quad \text{Card}(A) = 1 \rightarrow \text{Card}(\wp(A)) = 2$$

$$\wp(a, b) = \{\emptyset, \{a\}, \{b\}, \{a, b\}\} \quad \text{Card}(A) = 2 \rightarrow \text{Card}(\wp(A)) = 4$$

$$\text{Card}(A) = n \rightarrow \text{Card}(\wp(A)) = 2^n$$

Si se parte del conjunto vacío y de acuerdo con el axioma de pares se puede obtener un nuevo conjunto que contenga por elemento al conjunto inicial. Además, este axioma permite obtener otro conjunto cuyos elementos sean los dos anteriores y así sucesivamente, como se muestra a continuación:

$$0 = \emptyset$$

$$1 = \{\emptyset\}$$

$$2 = \{\emptyset, \{\emptyset\}\}$$

$$3 = \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}\}$$

$$\vdots$$

Con el anterior proceso de construcción se crea un conjunto infinito de manera potencial, pero este proceso es inacabado porque el problema radica en tomar todos los componentes, de tal forma que se pueda ver una colección que contengan los infinitos términos. Para dar solución a este problema se establece el axioma del infinito diciendo que existe un conjunto infinito que contiene al cero y a todos los elementos formados a partir de adicionar cada vez un nuevo componente.

La siguiente definición se presenta para formalizar el axioma del infinito.

Definición 1.4 *Un conjunto I es inductivo siempre que:*

I. $0 \in I$

II. Si $x \in I$, entonces el sucesor de x $S(x) \in I$

Con la anterior definición se puede formalizar el axioma del infinito:

Axioma Del Infinito: Existe un conjunto inductivo.

[9] Axioma de Regularidad o Fundamentación

Si A es un conjunto no vacío entonces A posee un elemento $\{a\} \in A$ tal que $\{a\} \cap A = \emptyset$.

$$\forall X (X \neq \emptyset \rightarrow \exists x \in X \wedge x \cap X = \emptyset)$$

Este axioma asegura que los conjuntos cumplen con la prohibición del regreso al infinito, es decir, garantiza que en esta teoría de conjuntos no suceden cadenas infinitas descendientes de pertenencias:

$$\dots, A_n \in A_{n-1} \in \dots A_3 \in A_2 \in A_1$$

Ni suceden cadenas circulares de pertenencias;

$$A_1 \in A_n \in A_{n-1} \in \dots A_3 \in A_2 \in A_1$$

Ni sucede que un conjunto sea elementos de sí mismo: $A \in A$.

Existen teorías axiomáticas consistentes que no incluyen el axioma de fundamentación y por tanto admiten conjuntos que son miembros de sí mismos. Ver en este sentido [1]. En este libro, Peter Aczel expone una teoría axiomática consistente que permite la existencia de conjuntos que forman cadenas de pertenencia con descenso infinito. En su propuesta incorpora el axioma de anti-fundamentación. En la teoría de conjuntos de Peter Aczel define a los conjuntos como objetos, que pueden ser grafos, y permite ver las

relaciones o propiedades que están inmersas en el conjunto.

Con el trabajo de Peter Aczel de conjuntos no bien fundados en [1], surgieron conceptos como el de *Coálgebras* que es dual a las álgebras, el de coinducción y corecursión que es dual a la inducción y recursión respectivamente. La dualidad de estos principios serán estudiados en los dos siguientes capítulos. Sin embargo, para el estudio de los números naturales desde la teoría de conjuntos en este capítulo se hará uso el Axioma de fundamentación y la teoría de ZFC, con el fin de comprender el proceso de construcción de los números naturales.

[10] Axioma de Reemplazo

Si $F(x, y)$ representa una condición en dos variables que es funcional en x , es decir, tal que para cualesquiera A, B, C colecciones, sin importar su tamaño ³

$$F(A, B) \wedge F(A, C) \rightarrow B = C.$$

El axioma dice que F representa una relación funcional, entonces para cada conjunto A su imagen por F es un conjunto.

[11] Axioma de Elección

Sea X un conjunto de conjuntos no vacíos. Una función de elección para X es una función $f : X \rightarrow \bigcup X$ tal que $((\forall a \in X) f(a) \in a)$

El axioma de elección es, tal vez, uno de los axiomas más controversiales en la historia de las matemáticas; aunque hoy en día es mayormente aceptado como cierto, el axioma de elección sigue siendo un tema controversial.

De acuerdo con [10], la teoría ZFC se “ construye ” partiendo de las siguientes premisas: todos los elementos de los conjuntos han de ser a su vez

³A, B, C pueden ser clases, es decir, sub-colecciones de la colección de todos los conjuntos.

conjuntos y dado que por el axioma de regularidad o fundamentación, no puede haber “ cadenas descendentes infinitas de pertenecía ”.

Los anteriores axiomas se presentaron con el fin de comprender la construcción tradicional de los números naturales, pero para presentar formalmente a $\mathbb{N}$, es necesario definir un conjunto que sea representante del cero, el cual es vacío, y evidentemente este conjunto no contiene elementos; luego es necesario definir los elementos que siguen del vacío, y para ello se emplea el concepto de sucesor.

1.3. El Sucesor

Definición 1.5 *Se llama sucesor de un conjunto X a $S(X) = X \cup \{X\}$*

Definición 1.6 *Si x es un conjunto, entonces $S(x)$ también lo es.*

1.4. Los Números Naturales

Intuitivamente se tiene que 0 representa la cantidad de elementos del conjunto $\emptyset$, pero, de su definición, el conjunto vacío tiene 0 elementos.

Definición 1.7 *Se definen los números naturales a través de las siguientes igualdades:*

$$0 = \emptyset$$

$$1 = S(0) = \{\emptyset\}$$

$$2 = S(1) = \{\emptyset, \{\emptyset\}\}$$

$$3 = S(2) = \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}\}$$

$$\vdots$$

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

$$\vdots$$

Es importante resaltar que:

- I. Cada una de las colecciones constituidas en la definición anterior es un conjunto.
- II. 0 es un número natural.
- III. Si n es un número natural, su sucesor también lo es.
- IV. Todo número natural n es 0 (cero) ó es el sucesor de algún otro número natural.
- V. La sucesión de números naturales se puede escribir como:

$$0 = \emptyset$$

$$1 = \{0\}$$

$$2 = \{0, 1\}$$

$$3 = \{0, 1, 2\}$$

$$\vdots$$

$$n + 1 = S(n) = \{0, 1, 2, 3, \dots, n\}$$

$$\vdots$$

Lo anterior asegura que a cada número natural le pertenece un número de la serie numérica $S(n)$, por ejemplo el número 3 es igual al conjunto $\{0, 1, 2\}$ y tiene como elementos los números 0, 1 y 2 que son los antecesores de 3. Pero si se quiere formar el conjunto de todos los números naturales, se requiere del axioma del infinito que asegura su existencia. Además con el axioma del infinito se puede concluir que los números naturales cumplen con la definición de conjunto inductivo, por lo que se puede establecer la siguiente observación:

Observación 1.5 *Dado un conjunto inductivo X , el conjunto*

$$\mathbb{N} = \{n \in X \mid n \text{ está en todos los conjuntos inductivos}\}$$

Es un conjunto inductivo. Además, si Y es otro conjunto inductivo, también se tiene

$$\mathbb{N} = \{n \in Y \mid n \text{ está en todos los conjuntos inductivos}\}$$

Con la observación 1.5 se puede concluir que cada número natural obligatoriamente debe pertenecer a todo conjunto inductivo, lo cual implica, que los números naturales son la intersección de todos los conjuntos inductivos, esto permite presentar la siguiente definición:

Definición 1.8 *El conjunto de los números naturales, designado por $\mathbb{N}$, es el conjunto:*

$$\mathbb{N} = \bigcap \{I : I \text{ es inductivo}\}$$

Hasta al momento se tiene que el conjunto de los números naturales pertenece a la colección todos los conjuntos inductivos y además, son la intersección de todos estos conjuntos, esto permite obtener la siguiente propiedad:

Observación 1.6 *$\mathbb{N}$ es el más pequeño de los conjuntos inductivos; es decir $\mathbb{N}$ es inductivo y para cualquier otro conjunto I , se tiene que $\mathbb{N} \subseteq I$.*

Hasta al momento con la definición 1.8 y la proposición 1.4, se puede concluir que los números naturales son el mínimo conjunto contenido en todos los conjuntos inductivos, **esto implica que $\mathbb{N}$ es un objeto minimal.**

Con el resultado anterior y de acuerdo con Ortíz y Valencia en [17], se establece que los números naturales son un objeto inicial, lo cual permite definir en ellos los principios de inducción y de recursión. La inducción es un principio de demostración y la recursión es un método que permite definir las propiedades de objetos estudiados intuitivamente. A continuación se estudiarán estos dos principios.

1.5. Principio de Inducción

En este punto ya se puede abordar el principio de inducción puesto que ya se reconoció la existencia de los números naturales. A grandes rasgos el principio de inducción va ligado a un razonamiento a través del cual se parte de una premisa dada para llegar a una conclusión, es decir, partir de lo particular a lo general. Este razonamiento se manifiesta en un método de demostración. Además, existen tres variantes del principio de inducción:

Principio de Inducción: Si un conjunto de naturales contiene al primer elemento, y si cada vez que contiene a un natural contiene también a su sucesor, entonces el conjunto contiene a todos los números naturales. Esto es, si $A \subseteq \mathbb{N}$ y si:

$$\text{I. } 0 \in A$$

$$\text{II. } n \in A \rightarrow s(n) \in A$$

$$\text{Entonces } A = \mathbb{N}$$

Demostración Como A es un conjunto inductivo entonces $\mathbb{N} \subset A$, por hipótesis se tiene que $A \subset \mathbb{N}$, lo que implica $A = \mathbb{N}$ ■

Principio de inducción matemática: Se denomina inducción a todo razonamiento que comprende el paso de proposiciones particulares a generales con la particularidad de que la validez de los últimos se deduce de la validez de los primeros. El método de inducción matemática es un método especial de demostración matemática que permite, a base de observaciones peculiares, juzgar de las regularidades generales correspondientes.

Teorema 1.1 Sea $\varphi(x)$ una fórmula de variable x que satisfice:

$$\text{I. } \varphi(0) \text{ se verifica}$$

$$\text{II. Para todo } n \in \mathbb{N}, \varphi(n) \text{ implica que } \varphi(n+1)$$

Entonces:

$$\varphi(n) \text{ se verifica para todo } n \in \mathbb{N}.$$

Demostración

De las hipótesis se deduce que $X = \{n \in \mathbb{N} / \varphi(n)\}$ es un conjunto inductivo. Por lo tanto, la proposición 1.3 implica que $\mathbb{N}$ está contenido en X , y por tanto $X = \mathbb{N}$ ■

A continuación se presenta el principio de Inducción Matemática Fuerte, este es un método de demostración matemática equivalente a la inducción matemática, pero más general.

Principio de Inducción Matemática Fuerte: Sea $A \subseteq \mathbb{N}$ tal que:

- I. $0 \in A$
- II. Si para todo $k \leq n$ y $k \in A$, entonces $n + 1 \in A$

Entonces:

$$A = \mathbb{N}$$

Demostración

Se define el conjunto:

$$B = \{n \in \mathbb{N} : k \in A \text{ para todo } k \leq n\}$$

Entonces $B \subseteq A$, y por la propiedad I se tiene que $0 \in B$. De otro lado, supongamos que $n \in B$, entonces por definición tenemos que $k \in A$ para todo $k \leq n$, y por la propiedad II se tiene que $n + 1 \in A$. Consecuentemente.

$$k \in A \text{ para todo } k \leq n + 1$$

Lo que significa $n + 1 \in B$. Por el principio de inducción se sigue que $B = \mathbb{N}$. Finalmente tenemos que $\mathbb{N} = B \subset A$, y por tanto $A = \mathbb{N}$. ■

El principio de inducción matemática fuerte o completa es equivalente al principio de inducción, solo que es más sencilla su aplicación en ciertos casos, y quizá más conveniente [16].

Hasta ahora se ha dicho que $\mathbb{N}$ es un conjunto inductivo y es el más pequeño de ellos, es decir, que $\mathbb{N}$ es un conjunto minimal. Esto implica que $\mathbb{N}$ es un objeto inicial y por ende se establece un principio de inducción y recursión [17].

Además, se puede establecer una relación de orden en $\mathbb{N}$ puesto que, la sucesión de los números naturales obtenidas en la definición 1.6 permite establecer la cadena $0 \in 1 \in 2 \in \dots$, evidenciando un orden en $\mathbb{N}$. Después de definir el orden en los naturales, se puede presentar el principio de recursión que garantiza que las operaciones en los números naturales estén bien definidas.

1.5.1. Orden en los Naturales

Definición 1.9 Si $n, m \in \mathbb{N}$ se dice que m es menor que n , que se simboliza por $m < n$, si y solo si, $m \in n$.

De acuerdo a lo anterior se tiene: $0 < 1 < 2 < 3 < \dots$

Ejemplo

$2 \leq 3$, porque $2 = \{\emptyset, \{\emptyset\}\}$,

$3 = \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}\}$ y $\{\emptyset, \{\emptyset\}\} \subseteq \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}\}$

Cuando se cumple $x \leq y$ decimos que x es menor o igual que y .

Con la anterior definición de orden se tiene el siguiente teorema:

Teorema 1.2 ⁴ Para todo $n, m \in \mathbb{N}$, se tiene que:

- I. $0 \leq n$
- II. Si $S(n) = S(m)$, entonces $n = m$
- III. Si $m \in n$, entonces $m \in \mathbb{N}$
- IV. $m < S(n)$ si y solo si $m \leq n$

⁴La demostración de este teorema se puede realizar por el principio de inducción, o también se puede encontrar la demostración en [16] en las páginas 58 y 59.

V. Si $m < n$, entonces $S(m) \leq n$

Teorema 1.3 $\leq$ es un orden total discreto sobre $\mathbb{N}$.

Demostración

- I. $\leq$ es reflexiva pues para todo $n \in \mathbb{N}$, $n \leq n$
- II. $\leq$ es antisimétrica pues si $(n \leq m \wedge m \leq n)$ y $m \neq n$, se tendría que $n \in m \wedge m \in n$ entonces $n \in n$ lo cual es una contradicción.
- III. $\leq$ es transitiva porque $(n \leq m \wedge m \leq s) \Rightarrow n \leq s$
 $(n \in m \wedge m \in s) \Rightarrow n \in s \Rightarrow n \leq s$

Por último veamos que es discreto; es decir, que $n < S(n)$ y que no existe un natural tal que $n < m < S(n)$. En primera instancia $n \in S(n)$, por tanto $n < S(n)$. Ahora si suponemos que existe un natural m tal que $n < m < S(n)$, entonces por el ítem III del teorema anterior se tiene que $m \leq n < m$, de donde se tendría que $m < m$, que es una contradicción. ■

Con el teorema anterior se demuestra que los números naturales poseen un orden total discreto. Ahora bien el concepto de sucesor se puede presentar como la función $S : K \rightarrow X$ en la cual cada elemento $n \in K$ le corresponde un $S(n) \in X$. La función sucesor respeta estrictamente el orden porque es estrictamente monótona en K , es decir, si $x, y \in K$ y $x < y \Rightarrow S(x) < S(y)$, esto es cierto porque: $x < y \Rightarrow S(x) \in S(y)$. Además, como $y \in S(y)$ y $S(x) \in y$ pues si no fuese así entonces $S(x) = y$, lo que implica que $S(x) \leq y < S(y)$

Hasta al momento se tiene que los números naturales poseen un orden total discreto y son el conjunto inductivo más pequeño, considerándose como un objeto minimal o inicial, esto permitió definir el principio de inducción como un método de demostración, sin embargo, hace falta el principio de recursión, el cual garantiza la buena definición de las operaciones suma y producto en $\mathbb{N}$. Esto se hace con el fin de presentar una estructura álgebraica de los números naturales.

1.5.2. Principio de Recursión

La recursividad es un método que aparece en la solución de muchos problemas, por ejemplo, en la definición de conjuntos, sucesiones, funciones o diseño de algoritmos.

De acuerdo con [5] el concepto de recursión generaliza el de función sucesor, dado que, en lugar de asignar a cada natural su sucesor se le asigna un elemento de un conjunto cualquiera.

La recursividad se caracteriza por la existencia de:

- I. Unas reglas básicas que dan definición para unos casos concretos.
- II. Unas reglas recursivas que dan la definición del problema mediante referencias a una versión más simple de este.

Los anteriores items permite establecer que un concepto se concreta recursivamente cuando: se define explícitamente para los primeros casos y se da una lista de reglas que lo define para el caso n -ésimo. La importancia de las definiciones recursivas radica en que se da un método constructivo para encontrar los términos de la sucesión.

Ejemplo⁵

El concepto de potenciación se puede definir recursivamente así:

Para $a, n \in \mathbb{N}$ definimos: $a^1 = a$, y $a^n = a^{n-1}a$ para todo $n \geq 2$; de esta manera tendríamos por ejemplo:

$$a^2 = a^{2-1}a = a^1a = aa, a^3 = a^{3-1}a = a^2a = aaa \text{ y así sucesivamente.}$$

De acuerdo con [20] se puede afirmar que toda definición recursiva siempre define una sucesión en un determinado conjunto X . Entonces la importancia de una buena definición recursiva es el principio de inducción. A continuación se formalizará el teorema de recursión de Dedekind tomado de [9], con el fin de garantizar que las operaciones suma y producto de $\mathbb{N}$ estén bien definidas.

⁵Este ejemplo es tomado de [20]

Teorema 1.4 (*Recursión de Dedekind*). Si A es un conjunto que contiene por lo menos un elemento $a \in A$ y $g : A \rightarrow A$ es una función de A en sí mismo, entonces existe una única función $\varphi : \mathbb{N} \rightarrow A$ que satisface las dos propiedades siguientes:

$$I. \quad \varphi(0) = a$$

$$II. \quad \varphi \circ S = g \circ \varphi$$

La función φ se dice que está definida recursivamente a partir de $\varphi(0) = a$ con formula de recursión $\varphi(n+1) = g(\varphi(n))$.

Demostración

Supongamos que dicha función existe y procedamos a demostrar primero su unicidad. Para ello sea φ_1, φ_2 dos funciones de $\mathbb{N}$ a A que satisfacen (1) y (2), se demostrara haciendo inducción sobre n que $\varphi_1(n) = \varphi_2(n)$ para todo n . La base inductiva se da por (1) pues $\varphi_1(0) = a = \varphi_2(0)$ y nuestra hipótesis de inducción estipula que para todo $K > 0$ se cumple que $\varphi_1(K) = \varphi_2(k)$. Es por (2) que vemos que se cumple el paso inductivo:

$$\varphi_1(K+1) = \varphi_1(S(k)) = g(\varphi_1(K)) = g(\varphi_2(K)) = \varphi_2(k+1)$$

Para demostrar la existencia de dicha φ considérense todos los subconjuntos $H \subset \mathbb{N} \times A$ que cumplen las dos propiedades dadas a continuación:

I. La pareja ordenada $(0, a) \in H$, y

II. Para todos n, b ; si se cumple que $(n, b) \in H$ entonces $(S(n), g(b)) \in H$.

Como todo el conjunto $\mathbb{N} \times A$ es uno de los mencionados subconjuntos H y todos los H contiene al elemento $(0, a)$, entonces la intersección de todos ellos -Se puede nombrar como D - es no vacía y es además el subconjunto más pequeño de $\mathbb{N} \times A$ que satisface (1) y (2) en el sentido de que para cualquier otro subconjunto $G \subset \mathbb{N} \times A$ que satisfaga ambas propiedades se tiene que $D \subseteq G$. Se afirma que el subconjunto D es la gráfica de una función $\varphi : \mathbb{N} \rightarrow A$ y se probará esta afirmación con el principio de inducción fuerte

con la siguiente propiedad:

(*) A cada $n \in \mathbb{N}$ le corresponde un y solo un $b \in A$ tal que $(n, b) \in D$. Para comenzar la inducción nótese que por (1): $(0, a) \in D$. Si existiese $(0, c) \in D$ tal que $c \neq a$ entonces se podría quitar $(0, c)$ de D y obtener un conjunto $D - \{(0, c)\}$ que todavía satisface (1) y (2); pero como esto contradice el que D es el más pequeño de los subconjuntos de $\mathbb{N} \times A$.

Para completar la demostración por inducción se plantea que para cada $k \in \mathbb{N}$ tal que $k > 0$ existe una única $b \in A$ tal que $(k, b) \in D$. Por (2) se tiene que $(S(k), g(b)) \in D$. Si se tuviese $(S(k), c) \in D$ con $c \neq g(b)$, se puede remover a $(S(k), c)$ de D y llegar a una contradicción por el mismo razonamiento utilizado en el párrafo anterior. De esta manera se termina la demostración por inducción y se establece la validez de (*) para todos los naturales. En consecuencia podemos escribir a D como la gráfica de una función $\varphi : \mathbb{N} \rightarrow A$ como sigue:

$$D = \{(n, \varphi(n)) | n \in \mathbb{N}\}$$

La propiedad (1) significa que $\varphi(0) = a$ y la (2) que $(S(k), g(\varphi(n))) \in D$, por lo que se concluye que $\varphi \circ S(n) = g \circ \varphi(n)$ para toda n . ■

1.5.3. Operaciones Con Números Naturales

A continuación se definen las operaciones de suma y producto de los números naturales, esto se hace con el fin de evidenciar una estructura algebraica en $\mathbb{N}$ y de acuerdo con Ortíz y Valencia en [17], el principio de recursión garantiza la buena definición de las funciones recursivas, entre las que se pueden rescatar las operaciones de suma y multiplicación en los números naturales, que constituyen parte fundamental en las presentaciones básicas de la estructura algebraica de $\mathbb{N}$.

Suma

La suma de dos naturales se define como el resultado de aplicar al primero la función sucesor iterada el número de veces que indica el segundo. En particular sumar la unidad equivale a aplicar la función sucesor (iterada una vez).

Se define la suma de los números naturales $n, m \in \mathbb{N}$, denotada $n + m$ como $n + m = +_n(m) \in \mathbb{N}$

Los números naturales dados se denominan sumandos, n es el primer sumando y m es el segundo sumando, también diremos que la suma $n + m$ es incrementar n en m .

Definición 1.10 $+_n : \mathbb{N} \rightarrow \mathbb{N}$

$$I. +_n(0) = n$$

$$II. +_n(s(k)) = s(+_n(k))$$

Producto

Multiplicar dos números naturales es iterar la suma del primero consigo mismo tantas veces como lo indique el segundo. En otras palabras, la multiplicación es una suma iterada siendo los sumandos todos iguales al primer número natural, y estando repetida la iteración las veces que indique el segundo. Al multiplicar puede haber ninguna iteración, una iteración, dos iteraciones, etc.

Se define el producto n por m , denotado nm o también $n \cdot m$, como

$$nm = n \cdot m = \bullet_n(m)$$

Definición 1.11 $\bullet_n : \mathbb{N} \rightarrow \mathbb{N}$

$$I. \bullet_n(0) = 0$$

$$II. \bullet_n(s(k)) = \bullet_n(k) + n$$

Observación 1.7 *La suma cumple con las propiedades: Asociativa, conmutativa y tiene elemento neutro. Por otro lado, el producto cumple las propiedades de: Elemento neutro, distributiva, conmutativa y asociativa.*⁶

Con las operaciones suma y producto definidas en los números naturales se puede establecer que este conjunto posee una estructura algebraica, además, es un conjunto bien ordenado, inductivo y minimal.

⁶Las demostraciones de estas propiedades se encuentran en [12] en las páginas 46-54

En este capítulo se han presentado los axiomas del sistema ZFC, con el fin de proveer un lenguaje conjuntista, que permita mostrar o identificar las operaciones y propiedades de los números naturales, y en un sentido muy especial reconocer el principio de inducción. Sin embargo, lo más importante de este capítulo es reconocer a los números naturales como el subconjunto de los conjuntos inductivos, que cumple con la característica de ser la intersección de todos los conjuntos inductivos, esto permite establecer que $\mathbb{N}$ es un conjunto minimal. En otras palabras como se menciona en [17] los números naturales son un objeto inicial y al ser inicial cumple con los principios de inducción y recursión, esto permite establecerlos como una estructura algebraica. En los próximos capítulos se volverá a este conjunto para hacer la dualidad con los números reales.

Capítulo 2

Inducción y Coinducción

En este capítulo se retomarán los principios de inducción y recursión estudiados anteriormente, con el fin de comprender los principios de coinducción y corecursión. En consecuencia se presentará una dualidad entre estos principios por medio de la teoría de puntos fijos y la teoría de retículos, con el objetivo de definir un conjunto inductivo y coinductivo empleando estas teorías. Estos resultados permiten estudiar a $\mathbb{N}$ y $\mathbb{R}$ por medio de la teoría de categorías.

Los conceptos relacionados con conjuntos inductivos y coinductivos se formalizan por medio de la teoría de puntos fijos, como lo menciona Cejas en [4]:

“Poincaré demostró que las soluciones de ciertos problemas analíticos se podían estudiar definiendo un conjunto X y una función $f : X \rightarrow X$ de tal forma que las soluciones del problema correspondiesen a los puntos fijos de la función f , esto es, los puntos $x \in X$ tales que $f(x) = x$ ”.

A continuación se formaliza la inducción y la coinducción mediante la teoría de puntos fijos, para luego abordar estos mismos principios desde la teoría de retículos.

2.1. Teoría de puntos fijos

La teoría de puntos fijos es parte integral de la topología. En esta teoría se establece un conjunto X y una función $f : X \rightarrow X$, la cual permite estudiar

y comprender los objetos inductivos y coinductivos.

De acuerdo con Cejas en [4], un punto fijo es un punto cuya imagen por una función es él mismo; es decir, si $x \in X$ entonces $f(x) = x$.

Se sabe por el axioma de extensión y la definición 1.1, que: $f(x) = x \leftrightarrow f(x) \subseteq x \wedge x \subseteq f(x)$, al simplificar lo anterior se tiene que: $f(x) \subseteq x$ y se conoce como punto prefijo y $x \subseteq f(x)$ es llamado punto postfijo. Más adelante se definirán estos dos conceptos, con el objetivo de identificar conjuntos inductivos y coinductivos. Para lograr esto, es necesario definir una topología y se parte de la siguiente afirmación.

Definición 2.1 (*Operador monótono:*) Sea $P(A)$ el conjunto potencia de un conjunto A , un operador sobre A es una función $F : P(A) \rightarrow P(A)$. Un operador es monótono si se cumple la siguiente implicación.

$$X \subseteq Y \subseteq A \text{ entonces } F(X) \subseteq F(Y)$$

Es decir, un operador o función es monótona si preserva el orden.

Con la definición anterior se tiene un conjunto potencia y un operador monótono. A continuación se va a establecer una topología partiendo del conjunto potencia, para esto se definen: espacio topológico, vecindad y adherencia. Esto se hace con el propósito de estudiar los puntos prefijos y postfijos, que ayudarán a evidenciar la dualidad existente entre los conjuntos inductivos y coinductivos.

Definición 2.2 (*Espacio topológico*) una topología sobre un conjunto X es una colección τ de subconjuntos de X con las siguientes propiedades:

- I. $\emptyset$ y X están en τ .
- II. la unión de los elementos de cualquier subcolección de τ está en τ .
- III. la intersección de los elementos de cualquier subcolección finita de τ está en τ .

La anterior definición es tomada de [15]. A continuación se define la noción vecindad y clausura de un conjunto en un espacio topológico, estas definiciones abrirán paso para formalizar el concepto de punto postfijo y prefijo.

Definición 2.3 (*Definición de vecindad:*) En un espacio topológico X , se llama vecindad ν de un subconjunto de A de X , a todo conjunto abierto que contiene a A .

Definición 2.4 (*Definición de adherencia o clausura:*) x es adherente o de clausura A , si toda vecindad de x tiene al menos un punto de A , es decir:

$$x \in \bar{A} \leftrightarrow \{\forall \nu_x \in \tau_A \text{ tal que } \nu_x \cap A \neq \emptyset\}$$

La clausura o adherencia de A se denota como $\mathcal{C}(A)$ o $\bar{A}$

Con los anteriores definiciones se observa que $\mathcal{C} : \wp(A) \rightarrow \wp(A)$ es un operador de clausura, es decir como $\mathcal{C}$ manda de partes en partes, entonces entonces el operador $\mathcal{C}$ manda al conjunto A en su adherencia, y se denota $\mathcal{C} : A \mapsto \bar{A}$.

Algunas de las más importantes propiedades de la clausura se resumen en el siguiente resultado.

$\mathcal{C}$ es un operador de clausura, es decir:

- I. $A \subset \mathcal{C}(A)$
- II. $\mathcal{C}\mathcal{C}(A) \subset \mathcal{C}(A)$
- III. $A \subset B \rightarrow \mathcal{C}(A) \subset \mathcal{C}(B)$

Definición 2.5 (*Conjunto cerrado*). Sea (X, τ) y sea $F \subset X$. El conjunto F se llama cerrado (respecto a la topología τ) si $F^c \in \tau$

Propiedad 2.1 (*Cerrado*): Sea X un espacio topológico y $A \subset X$, entonces A es cerrado, si y solo sí, $A = \bar{A}$.

Demostración

Si A es cerrado, $\bar{A} \subset A$, lo que implica que $A = \bar{A}$ y $x \notin A$, existe $V \in \mathcal{V}_x$ tal que $V \cap A = \emptyset$. Entonces $V \subset A^c$, de donde A^c es abierto y A es cerrado.

Como A es cerrado entonces cumple:

$$A \subset \bar{A} \text{ y } \bar{A} \subset A \blacksquare$$

La definición y propiedad de cerrado permiten comprender la definición de los puntos prefijos, esto se hace para poder identificar los conjuntos inductivos. De forma dual se definirá los puntos postfijos que abren paso para reconocer los conjuntos coinductivos.

Con lo anterior se puede enunciar la definición de F – Cerrado tomado de [7]

Definición 2.6 F -Cerrado o punto prefijo de F si $F(K) \subseteq K$

De los puntos prefijos en teoría de conjuntos se tiene que: $f(A) \subseteq A$, esto quiere decir, que las imágenes de $f(A)$ están contenidas en A , lo que implica que la intersección de todos los puntos prefijos es un conjunto mínimo, es decir:

$$\bigcap \{K \mid F(K) \subseteq K\}$$

Hasta ahora se ha presentado la definición de punto prefijo, por medio de la teoría de puntos fijos. A continuación dualmente se define los puntos postfijos. Esto se hace para definir el ínfimo y supremo de los puntos prefijos y postfijos respectivamente, estos resultados serán útiles para poder presentar un conjunto inductivo y coinductivo.

Definición 2.7 F es denso o Punto posfijo si $K \subseteq F(K)$:

Los puntos post-fijos en teoría de conjunto se tiene que $A \subseteq F(A)$ esto quiere decir que A está contenido en las imágenes $F(A)$, por lo tanto, la unión de todos los puntos postfijos son un conjunto máximo, es decir:

$$\bigcup \{K \mid K \subseteq F(K)\}$$

Con las definiciones punto prefijo y postfijo, se evidencia la existencia del mínimo y del máximo en los conjuntos, más adelante usando la teoría

de retículos se presentará el teorema de Knaster-Tarski que contiene dos definiciones equivalentes a las anteriores. Como uno de los objetivos de este capítulo es definir los conjuntos inductivos y coinductivos, se presenta la definición de ínfimo y supremo de los puntos prefijo y postfijo respectivamente, que hace uso de las definiciones 2.5 y 2.6.

La siguiente definicion es tomada de [7]

Definición 2.8 (*Ínfimo y supremo de los puntos prefijos y postfijos*). Se definen los siguientes conjuntos:

- El ínfimo de los puntos prefijos de F es: $\bigcap \{K \mid F(K) \subseteq K\}$
- El supremo de los puntos postfijos de F es: $\bigcup \{K \mid K \subseteq F(K)\}$

Hasta al momento se tiene que la intersección de los puntos prefijos son el ínfimo y la unión de los puntos postfijos son el supremo, falta mostrar la relación que existe entre los conjuntos inductivos y coinductivos con los puntos prefijos y posfijos respectivamente. Esta relación se abordará por medio de la teoría de retículos y el teorema de Knaster - Tarski, que permite presentar una dualidad entre los conjuntos inductivos y coinductivos, con el fin de evidenciar la dualidad que existe entre $\mathbb{N}$ y $\mathbb{R}$.

2.2. Retículos

A continuación se presentarán algunos conceptos de la teoría de retículos. Sin embargo, el objetivo de este trabajo no es estudiar exhaustivamente esta teoría. Un estudio detallado de la teoría de retículos puede verse, por ejemplo, en [2] ó [21].

Un *retículo* es una terna $(L, \vee, \wedge)$ donde $L \neq \emptyset$ es un conjunto y $\vee, \wedge$ son dos operaciones binarias en L verificando las propiedades: sea $a, b, c \in L$

I. Asociativa:

$$(a \vee b) \vee c = a \vee (b \vee c)$$

$$(a \wedge b) \wedge c = a \wedge (b \wedge c)$$

II. Conmutativa:

$$a \vee b = b \vee a$$

$$a \wedge b = b \wedge a$$

III. Absorción:

$$a \vee (a \wedge b) = a$$

$$a \wedge (a \vee b) = a$$

La propiedad de idempotencia $a \vee a = a$ y $a \wedge a = a$ se puede demostrar con las tres anteriores [21].

Las anteriores propiedades hacen parte de la definición algebraica de los retículos. Sin embargo, los retículos también se definen por medio del orden; es decir, están dotados de un orden que cumple con las propiedades: reflexiva, antisimétrica y transitiva, además pueden tener supremo, ínfimo o ambos.

A continuación se presentarán dos ejemplos de retículo:

Ejemplo

El conjunto ordenado $(N, |)$ es un retículo. En este caso se tiene que $x \vee y = mcm(x, y)$ mientras que $x \wedge y = mcd(x, y)$. De la misma forma, si $n \in N$ entonces $D(n)$, con el orden dado por la divisibilidad es un retículo. Supremo e ínfimo vienen dados por el mínimo común múltiplo y el máximo común divisor respectivamente.

Ejemplo

Dado X un conjunto, $(\wp(X), \cap, \cup)$ es un retículo.

Proposición 2.1 *Sea L un conjunto que tiene definidas las operaciones $\vee$ y $\wedge$ que satisfacen las propiedades: conmutativa, asociativa, idempotencia y*

la absorción. Supongamos que en L definimos la relación:

$$x \leq y \text{ si } x \vee y = y$$

Entonces, $(L, \leq)$ es un conjunto parcialmente ordenado tal que todo conjunto infinito tiene supremo e ínfimo.

La siguiente demostración se realizará por medio de la definición de orden pero si se desea estudiar la demostración por medio de la definición algebraica se puede encontrar en [8].

Demostración

Se debe mirar si tiene $(L, \leq)$, ínfimo y supremo.

I. Veamos en primer lugar que $(L, \leq)$ es un conjunto ordenado. Para esto, comprobemos que la relación $\leq$ es:

- Reflexiva: como $x \vee x$ se tiene que $x \leq x$ para cualquier $x \in L$.
- Antisimétrica: supongamos que $x \leq y$ e $y \leq x$. Esto implica que $x \vee y = y$ y $y \vee x = x$. Puesto que $\vee$ es conmutativa deducimos que $x = y = (x \vee y)$.
- Transitiva: supongamos ahora que $x \leq y$ y que $y \leq z$, es decir, $x \vee y = y$ e $y \vee z = z$. Entonces:

$$x \vee z = x \vee (y \vee z) = (x \vee y) \vee z = y \vee z = z$$

luego $x \leq z$.

II. Tiene supremo: dados $x, y \in L$ se verifica que $\sup(\{x, y\}) = x \vee y$. Puesto que $x \vee (x \vee y) = (x \vee x) \vee y = x \vee y$ se tiene que $x \leq x \vee y$. De la misma forma se comprueba que $y \leq x \vee y$.

Si $x \leq u$ e $y \leq u$ (es decir, $x \vee u = u$ (por hipótesis) e $y \vee u = u$). Entonces:

$(x \vee y) \vee u = x \vee (y \vee u) = x \vee u = u$ de donde se concluye $(x \vee y) \leq u$.

III. Tiene ínfimo: $\inf(\{x, y\}) = x \wedge y$ $(x \wedge y) \vee x = x \vee (x \wedge y) = x$. Luego $x \wedge y \leq x$. De la misma forma se comprueba que $x \wedge y \leq y$. Si $u \leq x$ y $u \leq y$ (es decir, $u \vee x = x$ y $u \vee y = y$) se tiene que:

$$U \wedge x = u \wedge (u \vee x) = u$$

$$u \wedge y = u \wedge (u \vee y) = u$$

$$u \wedge (x \wedge y) = (u \wedge x) \wedge y = u \wedge y = u$$

$$u \vee (x \wedge y) = (u \vee (x \wedge y)) \vee (x \wedge y) = (x \wedge y) \vee ((x \wedge y) \wedge u) = x \wedge y$$

Luego $u \leq x \wedge y$. ■

Observación 2.1 *Un retículo es un conjunto ordenado $(L, \leq)$ en el que todo conjunto finito tiene supremo e ínfimo, entonces:*

$x \vee y = \sup\{x, y\}$ y $x \wedge y = \inf\{x, y\}$ por lo tanto $(L, \vee, \wedge)$ es un retículo.

Si $(L, \leq)$ es un retículo y L tiene máximo, se denotará a éste por 1, mientras que si tiene mínimo se denotará por 0.

Por tanto:

$$\text{Máximo: } x \vee 1 = 1, x \wedge 1 = x$$

$$\text{Mínimo: } x \vee 0 = x, x \wedge 0 = 0$$

Es conveniente conocer que los retículos infinitos no siempre tiene máximo y mínimo, por ejemplo, $(\mathbb{N}, \leq)$ tiene mínimo pero no tiene máximo; $(\mathbb{Z}, \leq)$ no tiene ni mínimo ni máximo.

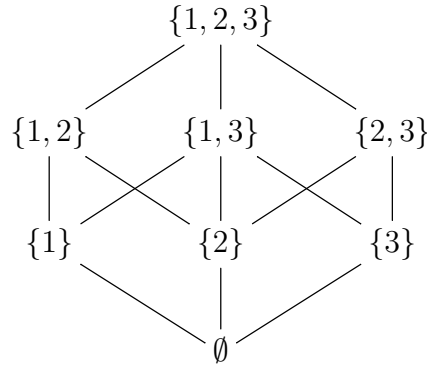
Como ya se revisó la definición de retículo entonces ahora se presentará la definición de retículo completo con el fin de presentar el teorema de Knaster-Tarski el cual permite evidenciar la dualidad que existe entre los conjuntos inductivos y coinductivos.

Retículo Completo

$(L, \leq)$ es un *retículo completo* si para todo $S \subseteq L$ existe $\bigwedge S$ (mínimo) ó $\bigvee S$ (máximo).

Ejemplo

Sea $A = \{1, 2, 3\}$ y $(P(A), \subseteq)$ un retículo completo.



Se puede observar que el máximo es el conjunto A y el mínimo es el conjunto $\emptyset$.

El siguiente teorema es tomado de [21], con el fin de establecer la existencia de los puntos fijos; es decir, que todo operador monótono tiene un punto fijo inductivo y un punto fijo coinductivo.

Teorema 2.1 (*PUNTO FIJO DE KNASTER-TARSKI*): Sea L un retículo completo y $f : L \rightarrow L$ una función monótona. Entonces el conjunto de puntos fijos de f en L es también un retículo completo.

Para comprender lo anterior, [21] enuncia el mismo teorema en otros términos que permiten evidenciar la existencia del menor punto prefijo y el mayor punto postfijo de f :

Sea L un retículo completo y $f : L \rightarrow L$ una función monótona. Entonces:

$$\bigwedge \{x \in L : f(x) \leq x\}.$$

Es el menor punto prefijo de f . Dualmente, f tiene el mayor punto post-fijo, dado por

$$\bigvee \{x \in L : f(x) \geq x\}.$$

Como los retículos no pueden ser vacíos, el teorema en particular garantiza la existencia de por lo menos un punto fijo. Esta afirmación es tomada de [22] que propone el teorema de Knaster-Tarski de la siguiente manera:

Teorema 2.2 *Sea L un retículo completo y $f : L \rightarrow L$ una función monótona. Entonces f tiene un punto fijo.*

Demostración

Sea $H = \{x \in L / f(x) \leq x\}$ y a el ínfimo de H . Vamos a demostrar que a es un punto fijo de f ; es decir, que $f(a) = a$. Para ello, se demostrará que $f(a) \leq a$ y que $a \leq f(a)$.

Para demostrar que $f(a) \leq a$, por la definición de a , basta probar que $f(a)$ es una cota inferior de H . En efecto, sea $x \in H$. Entonces:

- $a \leq x$ (por la definición de a)
- $f(a) \leq f(x)$ (porque f es monótona)
- $f(x) \leq x$ (porque $x \in H$).

Por tanto, $f(a)$ es una cota inferior de H .

Para demostrar que $a \leq f(a)$, por la definición de a , basta probar que $f(a) \in H$. En efecto, tenemos que:

- $f(a) \leq a$ (por el apartado anterior)
- $f(f(a)) \leq f(a)$ (por ser f monótona)

Por lo tanto, usando la definición de H , tenemos que $f(a) \in H$ ■

Con los anteriores resultados de la teoría de puntos fijos y retículos, se puede definir un conjunto inductivo y coinductivo, dado que, en las dos teorías se concluye la existencia del ínfimo μ y la existencia del supremo ν ;

$$\mu F = \bigwedge \{x \in L : f(x) \leq x\} \text{ y}$$

$$\nu F = \bigvee \{x \in L : f(x) \geq x\}.$$

Con esto se puede definir los conjuntos inductivos y coinductivos de la siguiente forma:

Conjunto Inductivo

$$\mu F = \bigcap \{X \subseteq L : F(X) \subseteq X\}$$

Conjunto Coinductivo

$$\nu F = \bigcup \{X \subseteq L : X \subseteq F(X)\}$$

Lo anterior evidencia que los conjuntos inductivos y coinductivos son duales, porque el primero es la intersección de todos los puntos prefijos mientras que el segundo es la unión de todos los puntos postfijos. Hasta este punto y teniendo en cuenta lo expuesto en el capítulo uno los números naturales son un conjunto inductivo, y son un conjunto minimal tal como lo demostró Dedekind [6], sin embargo, aún falta caracterizar a los números reales como un

conjunto coinductivo, esto se hace con el fin de establecer la dualidad entre $\mathbb{N}$ y $\mathbb{R}$.

Como ya se han definido los conjuntos inductivos, coinductivos y el principio de inducción; entonces, ahora se puede definir el principio de coinducción.

2.3. Principio De Coinducción

Los números naturales se establecen como un objeto inicial según Ortíz y Valencia en [17], lo cual permite definir en ellos los principios de inducción y de recursión. La inducción es un principio de demostración de propiedades de los números naturales y la recursión garantiza la buena definición de las propiedades de objetos estudiados inductivamente. En el caso de los números naturales se define recursivamente así: 0 pertenece a los números naturales, y si n pertenece a $\mathbb{N}$ entonces el sucesor de n también. Estos principios permiten construir un objeto a partir de un elemento distinguido que es el cero o el conjunto vacío y un elemento cualquiera, a esto Pavlović y Pratt en [19] denominan aritmética.

Por otro lado, se pueden definir los números reales en forma dual al conjunto anterior. Es decir, los números reales son un objeto final como los define Hilbert en 1900: un cuerpo maximal arquimediano totalmente ordenado. Lo que implica que es un objeto final tal como lo muestran Ortíz y Valencia en [17] y por medio de la dualidad que existe entre los objetos iniciales y finales se pueden establecer los principios de coinducción y de corecursión en $\mathbb{R}$.

La coinducción es un término que se utiliza en el lenguaje de la computación y pertenece a las matemáticas contemporáneas, pero es poco conocida inclusive en el mundo de las matemáticas y en la educación. Recientemente en la comunidad matemática fue considerada como un principio lógico verdadero. Este principio se halla oculto en cada análisis matemático que es realizado por alguna persona o máquina. Es decir, la coinducción se genera en el análisis matemático, mientras que su dual la inducción se genera en la aritmética.

“El paso básico al infinito en el cálculo elemental es coinductivo, dual al paso inductivo al infinito en la aritmética elemental¹[19]”.

El proceso de coinducción en términos generales es partir de lo general a un primer elemento y la corecursión es dual a la recursión y por tanto permite establecer que las operaciones entre listas infinitas estén bien formadas. Para aplicar el principio de coinducción se requiere de objetos finales, los cuales pueden ser de naturaleza infinita u objetos infinitos no bien fundados; es decir, que no tengan el primer elemento. Un ejemplo de un tipo coinductivo son las listas de secuencias infinitas.

La coinducción es un principio que permite definir estructuras infinitas a partir de sus partes finitas. Para construir objetos usamos la recursión y construimos conjuntos inductivos; al dualizar, y partir de un conjunto infinito es necesario deconstruir² sus partes para obtener sus elementos, pero como el conjunto es infinito entonces este procedimiento nunca terminaría, por lo que se hace necesario encontrar una forma de agrupar estos infinitos elementos de tal forma que hagan el procedimiento terminable en algún momento. Al observar este conjunto infinito, se pueden identificar comportamientos comunes entre los elementos, lo cual implica que se puede agrupar de cierta manera estos infinitos elementos y así poder reconocer todos los elementos del conjunto; por ejemplo, cuando se tiene el conjunto de los números reales que es infinito y se desea conocer sus elementos, no se estudia número por número, ya que lo haría un proceso interminable, si no, que se estudian los comportamientos comunes de las partes del conjunto; es decir, más formalmente se busca conocer la estructura³ de $\mathbb{R}$ la cual permite agrupar de cierta manera todos los elementos de $\mathbb{R}$ y estudiar un conjunto infinito a partir de sus partes finitas que es la estructura.

El proceso que permite agrupar los elementos de un conjunto infinito es a través de una función, llamada función deconstructora (o de observación)

¹The basic passage to infinity in elementary calculus is coinductive, dual to the inductive passage to infinity in elementary arithmetic”.

²Se usará esta palabra que ya se ha utilizado en la literatura, porque no es exactamente destruir sino desbaratar para entender.

³El conjunto de números reales tiene varias estructuras: de orden, algebraica, topológica, etc.

que se enfocan en obtener las partes del conjunto infinito mediante su estructura. Estas funciones permiten obtener un sistema de conceptos coherentes enlazados, cuyo objetivo es precisar la esencia del objeto de estudio.

De acuerdo con [7], este proceso de deconstruir al conjunto coinductivo se emplea en la computación porque en ella se trabajan con listas infinitas, para programar, y se hace por medio de dos operadores *head* y *tail*, a partir de una lista infinita de elementos de un conjunto A podemos obtener la parte inicial de la lista, es decir la cabeza y observarla, y al quitar la cabeza del resto de la lista nos quedamos con la cola, que sigue siendo una lista infinita. Estas dos funciones destructoras se aplican para poder obtener los elementos de la lista *stream* y así observar una lista finita.

Ejemplos de coinducción de forma intuitiva

Supongamos que se tienen dos máquinas de chocolates de distintas empresas, cada una al insertar una moneda de x denominación da como transacción un chocolate a cambio. Como el resultado de ambas máquinas es un producto de las mismas características, entonces el principio de coinducción permite establecer que las dos máquinas de cierta manera son iguales, porque se obtiene un mismo resultado a pesar de no conocer su funcionamiento interno.

Otro ejemplo de coinducción es cuando una empresa, quiere saber el procedimiento interno de un producto de la competencia. Supongamos que este producto es un computador entonces la empresa que quiere conocer su funcionamiento, lo que hace, es desarmar el computador y observar cada pieza y su funcionamiento.

Se puede concluir que la teoría de puntos fijos permite formalizar el concepto de coinducción, que es dual a la inducción. Si al plantear un conjunto potencia $\wp(A)$ de un conjunto A , implica que un conjunto operador sobre A es una función $F : \wp(A) \rightarrow \wp(A)$ y si es monótona preserva el orden. Si F es cerrado implica que $\wp(A)$ es un subconjunto A , lo cual es un punto prefijo, y hace referencia a la inducción; y si F es denso entonces A es un subconjunto de $\wp(A)$ lo cual es un punto postfijo y hace referencia a la coinducción. En un espacio topológico X , por definición se tiene que $\emptyset$ y X pertenecen a la topología y también pertenecen la unión e intersección de los subconjuntos de

la topología; esta definición permite que los puntos prefijos y puntos postfijos estén dentro de una topología. Con el teorema de Knaster-Tarski se tiene que la intersección genera un ínfimo y la unión genera un supremo, el primero permite definir conjuntos inductivos mientras que el segundo por ser el dual define conjuntos coinductivos.

El principio de coinducción es poco conocido, pero es trabajado desde el campo de la computación para describir listas infinitas, estas listas tienen el mismo comportamiento que los números reales. La coinducción es un método implementado en el análisis como paso básico para el desarrollo del cálculo elemental; además, la complejidad matemática de los objetos básicos como los números reales, las funciones y las sucesiones, se van construyendo cuando se empieza la enseñanza del análisis. Entonces la dualidad que existe entre la inducción y la coinducción, se debe a la existencia de objetos iniciales y objetos finales. El primero son los naturales y el segundo son los números reales; para poder entender estos objetos se debe estudiar coálgebras y teoría de categorías, la cual permite una mirada más general y estructural de lo que nos ofrece la teoría de conjuntos, topología y retículos.

En el primer capítulo se logró definir a los números naturales como un conjunto inductivo y minimal; es decir, un objeto inicial que permite establecer los principios de inducción y recursión. En el segundo capítulo se alcanzó presentar una dualidad, entre, los conjuntos inductivos y coinductivos, por medio de la teoría de puntos fijos y retículos. Además, se logró a grandes rasgos identificar a los números reales como un conjunto coinductivo y maximal, es decir como, un objeto final.

Al utilizar la teoría de retículos la noción de minimalidad lleva inmersa la noción de inducción y recursión, y la noción de finalidad lleva inmersa la noción de coinducción y corecursión, lo cual implica que $\mathbb{R}$ es un objeto final y por ende se pueden establecer los principios de coinducción y corecursión. Sin embargo, la teoría de retículos de alguna manera es un lenguaje débil dentro de la matemática del siglo XX, porque las matemáticas de este siglo tienen un lenguaje más fuerte que se caracteriza por la noción de morfismos, como lo es la teoría de categorías; que se presentará en el siguiente capítulo y la cual permitirá reconocer la dualidad que existe entre $\mathbb{N}$ y $\mathbb{R}$.

Capítulo 3

Dualidad Entre $\mathbb{N}$ y $\mathbb{R}$ en Categorías

En el capítulo anterior se evidenció por medio de Retículos que la noción de minimalidad lleva inmersos los principios de inducción y recursión, y dualmente la noción de finalidad lleva inmerso los principios de coinducción y corecursión. Sin embargo, este mismo estudio se llevará a cabo en la teoría de categorías, que es un lenguaje más fuerte que permite estudiar la estructura de los objetos. Para cumplir este fin se paga una cuota de formalismo, esto implica introducir un poco del lenguaje de categorías, aunque no se va hacer una investigación exhaustiva en ésta teoría debido a la naturaleza del trabajo, lo que implica que se realizará un ligero acercamiento con algunos contenidos, puesto que se considera que no es necesario saber todos los conceptos de la teoría de categorías para saldar lo que se ha hecho con la teoría de retículos.

La dualidad se puede pensar como el reflejo de una persona en un espejo, como sucede en la lógica proposicional que por cada expresión valida hay otra, que obedece a la misma estructura. Partiendo de esto, el propósito de este capítulo es usar el lenguaje de categorías para evidenciar la dualidad que existe entre los números naturales y los reales. Para cumplir este propósito, se presenta a los números naturales como un álgebra inicial y de forma dual se presenta a los números reales como una coálgebra final. Al caracterizar a $\mathbb{N}$ como un objeto inicial implica que cumple con los principios de inducción y recursión. Y al caracterizar a $\mathbb{R}$ como un objeto final implica que cumple con los principios de coinducción y corecursión.

3.1. Categorías

La teoría de categorías es un lenguaje usado en la matemática del siglo XX, y de acuerdo con [3] es la generalización de la teoría de retículos, porque cuando se tiene un retículo o conjunto ordenado se puede observar a las categorías como la generalidad del orden. De igual forma sucede con: funtor-función monótona, álgebras- puntos prefijos, y coálgebras-postfijos. A continuación se presentará la definición de categoría y después con un ejemplo se evidenciará que las categorías son la generalización de los retículos.

Definición 3.1 Una categoría C está compuesta por:

- I. Una colección C_0 de objetos : $A, B, C, \dots, a, b, c, \dots$
- II. Una colección C_1 de flechas, morfismos, homomorfismos: $f, g, h, \dots$
- III. Toda flecha f tiene asociado un origen o dominio $\text{dom}(f)$ y un destino o codominio $\text{cod}(f)$. Escribimos $f : A \longrightarrow B$ ó $A \xrightarrow{f} B$ para indicar que $\text{dom}(f) = A$ y $\text{cod}(f) = B$.
- IV. Si $f : A \longrightarrow B$ y $g : B \longrightarrow C$ hay una flecha $g \circ f : A \longrightarrow C$ llamada la composición de f y g . La composición de flechas es asociativa.
- V. Para todo objeto A existe una flecha $1_A : A \longrightarrow A$ llamada la flecha identidad de A . Esta flecha es neutra para la composición.

Ejemplo

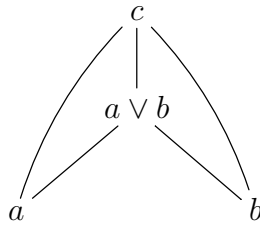
Sea un retículo o un conjunto parcialmente ordenado $\langle A, \leq \rangle$, entonces, se puede plantear la siguiente categoría C que cumple con:

- Los objetos son elementos de A
- Una flecha que en el retículo relaciona el orden de $a \leq b$ como $a \rightarrow b$. En la categoría, esta flecha se puede ver como un morfismo $f : a \rightarrow b$, tal que, $a \leq b$
- Existe una flecha identidad $1_a : a \rightarrow a$ tal que $a = a$

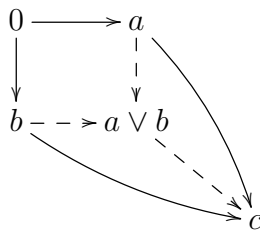
- La composición de flechas es asociativa; es decir, si $a \leq b \leq c$ entonces existen las funciones $f : a \rightarrow b$ y $g : b \rightarrow c$ lo que implica que hay una composición de flechas $g \circ f : a \rightarrow c$.

Hasta al momento el ejemplo permite evidenciar que las categorías son una generalización de los retículos puesto que cumple con la definición. Sin embargo, es importante conectar los resultados de la teoría de retículos con la teoría de categorías, es decir, en el segundo capítulo se obtuvo la existencia de un ínfimo y un supremo, ahora bien siguiendo el ejemplo las categorías también presentan estos dos conceptos, veamos como:

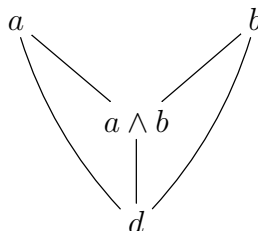
Si $\langle A, \leq \rangle$ es un retículo se tiene que para cada dos elementos existe un supremo c , que es $a \vee b$, gráficamente se puede observar lo siguiente:



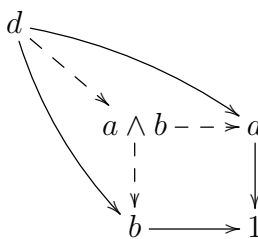
Como este retículo tiene un orden entonces se puede representar en categorías, pero es necesario retomar el resultado del capítulo 2 donde se representó al ínfimo como 0.



Por otro lado, en la definición de retículo se tiene que para cada dos elementos existe un ínfimo d , el cual es $a \wedge b$, esto se puede observar en el siguiente gráfico:



El anterior retículo se puede visualizar en la teoría de categorías y retomando que el 1 es la representación del supremo, entonces el conjunto ordenado quedaría de la siguiente forma:



En conclusión la teoría de categorías generaliza la teoría de retículos; esto permite retomar los resultados obtenidos en los anteriores capítulos y observarlos desde el lenguaje de categorías, permitiendo así presentar a los números naturales como un objeto inicial y a los números reales como un objeto final.

El ejemplo anterior fue una pequeña introducción de teoría de categorías por medio de los retículos, pero como se mencionó anteriormente la teoría de retículos es un lenguaje débil dentro de la matemática del siglo XX puesto que el lenguaje utilizado en esta época se caracterizaba por la noción de morfismo, por esta razón se definirán las categorías mediante los morfismos.

Para comprender la definición de categoría se requiere el concepto de homomorfismo, dado que las categorías trabajan con estas funciones.

Un homomorfismo, (o morfismo) desde un objeto matemático a otro de la misma categoría, es una función que preserva la estructura entre dos estructuras matemáticas relevantes. Por ejemplo, en un homomorfismo de orden, si

un objeto consiste en un conjunto X con un orden u y el otro objeto consiste en un conjunto Y con orden v , entonces:

$$f : X \Rightarrow Y \text{ satisface.}$$

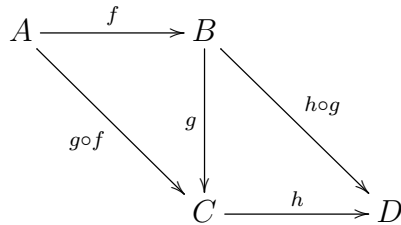
$$u \leq v \Rightarrow f(u) \leq f(v)$$

Por otro lado un morfismo $f : X \rightarrow X$ se denomina *endomorfismo* de X .

El concepto de morfismo se presentó con el fin de dar claridad a la definición 3.1 de categorías, es por ello que se intentará explicar el punto tres y cuatro.

Veamos una explicación del tercer punto (composición asociativa):

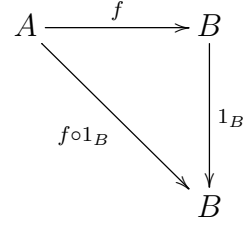
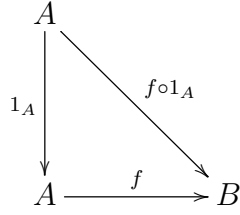
Sean los siguientes morfismos $f : A \longrightarrow B$, $g : B \longrightarrow C$ y $h : C \longrightarrow D$. Para que la composición sea asociativa debe cumplir la siguiente igualdad: $h \circ (g \circ f) = (h \circ g) \circ f$, para que esto se cumpla el siguiente diagrama debe conmutar¹:



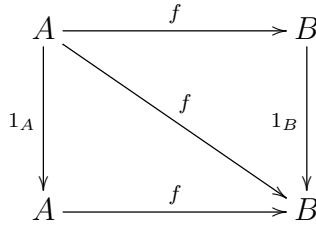
El diagrama conmuta si cumple esta igualdad $h \circ (g \circ f) = (h \circ g) \circ f$

Veamos una explicación del cuarto punto (identidad): un endomorfismo $f : X \longrightarrow X$, al cual se le llama *identidad*, es el elemento neutro de una composición y se denota 1_x . Si $f : A \longrightarrow B$, $1_A : A \longrightarrow A$ y $1_B : B \longrightarrow B$ son morfismos que cumplen las igualdades $f \circ 1_A = f$ y $1_B \circ f = f$.

¹Cuando un diagrama conmuta significa que para todo par de objetos a y b del diagrama, las composiciones que se hagan a lo largo de cualquier camino dirigido de a a b dan idéntico resultado.



$$f \circ 1_A = f \circ 1_B = f$$



A continuación se introduce el concepto de *funtor* que es uno de los más importantes de la teoría de categorías, se recomienda al lector consultar [11] desde las páginas 71 hasta la 80, dado que este concepto es muy general.

Ejemplo

Un conjunto parcialmente ordenado $(P, \leq)$ es una categoría pequeña, cuyos objetos son los miembros de P , y un morfismo de x a y ocurre cuando $x \leq y$.

el concepto de funtor, es fundamental para tener un criterio de comparación, entre las álgebras iniciales y las álgebras arbitrarias. De igual forma, con las coálgebras arbitrarias y las coálgebras finales.

Definición 3.2 *Un funtor es una función de una categoría a otra que lleva objetos a objetos y morfismos a morfismos de manera que la composición de morfismos y las identidades se preservan.*

A continuación se presentará un ejemplo tomado de [11] donde los funtores se denotan como $(\mapsto)$, sin embargo, después de este ejemplo se denotarán

con una flecha ($\rightarrow$).

Sea C una categoría y $A \in C$, $A \neq \emptyset$. Definimos un funtor $F : C \rightarrow C$. Si $B \in C$, $F(B) = B^A$ donde $B^A = \{f/f : A \rightarrow B\}$, el conjunto de morfismos de A en B , Considerado como grupo, por medio de la operación de grupo en B .

Si $f : B \rightarrow H$ es un homomorfismo, $f^A : B^A \rightarrow H^A$ es la función $f^A(g) = f \circ g$ para $g \in B^A$. Para describir el funtor, tenemos el siguiente diagrama.

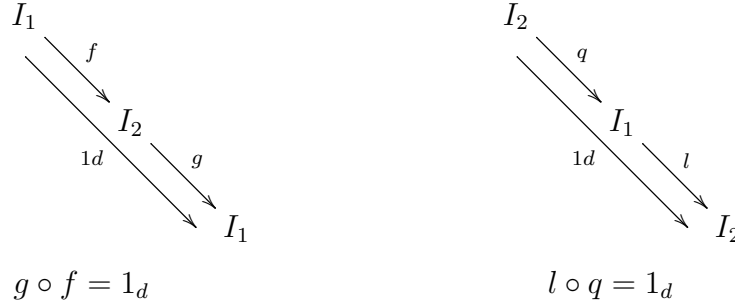
$$\begin{array}{ccc} B & \mapsto & B^A \\ f \downarrow & & \downarrow f^A \\ H & \mapsto & H^A \end{array}$$

En este ejemplo la flecha $\mapsto$ indica el valor del funtor sobre los objetos de B y H , entonces el ejemplo anterior indica que f^A es un funtor. El ejemplo evidencia que los funtores son esencialmente morfismos de categorías, es por esto que los funtores se hará referencia a morfismos únicos en lo que sigue del trabajo.

Como el fin de este trabajo es presentar la dualidad existente entre $\mathbb{N}$ y $\mathbb{R}$, se procede a presentar la definición de un objeto inicial y dualmente la de un objeto final desde la teoría de categorías.

Definición 3.3 : *Un objeto inicial de una categoría C es un objeto I en C tal que para todo objeto X en C existe un único morfismo $I \rightarrow X$. La noción dual es la de objeto final es decir, un objeto F es final si para todo objeto X en C existe un único morfismo $X \rightarrow F$.*

En una categoría pueden no existir objetos iniciales o finales, y si existen, son únicos, por ejemplo si I_1 y I_2 son dos objetos iniciales, entonces hay un único isomorfismo entre ellos. Además, si I es un objeto inicial, entonces cualquier objeto isomorfo a I es inicial. Por dualidad, todo lo anterior es cierto para objetos finales.



La Definición 3.3 es de suma importancia porque ayudará a identificar la dualidad de $\mathbb{N}$ y $\mathbb{R}$. Además, garantiza la existencia y unicidad de morfismo único que en el transfondo es un funtor. Lo que sigue en el trabajo se basará en clasificar a los números naturales como un objeto inicial y de forma dual a los números reales como objeto final. Esto se hará introduciendo álgebras iniciales y coálgebras finales.

En el capítulo 1 se estudiaron los números naturales y se planteó que eran una estructura algebraica y un conjunto minimal, es decir, que son un objeto inicial y que cumplen con los principios de inducción y recursión, ahora bien, se identificará a $\mathbb{N}$ como parte de las álgebras de tipo 0 – 1 con el fin de clasificarlos como álgebras iniciales para verlos desde una perspectiva categórica que permita evidenciar la dualidad que existe entre este objeto con los números reales.

3.2. Álgebras Tipo 0 - 1

En los capítulos 1 y 2 se presentó a los números naturales como un conjunto minimal, dicha minimalidad permite reconocer a $\mathbb{N}$ como un objeto inicial. Sin embargo en estos capítulos no se contaba con el lenguaje de categorías y por tanto no se había definiendo un objeto inicial desde esta teoría.

A continuación, Se formará la categoría C_{0-1} de las álgebras de tipo 0 – 1 ; es decir, $C_{0-1} = \{A : A \text{ es álgebra tipo } 0 - 1\}$ en la cual $A = \langle A, \oplus, a^0 \rangle$, con los morfismos $0 : 1 \rightarrow \mathbb{N}$ y $S : \mathbb{N} \rightarrow \mathbb{N}$. El primer morfismo es una relación 0-aria, mientras que el segundo es una relación 1-aria. De acuerdo con [17] los números naturales son un objeto dotado de una estructura, es decir, como

un álgebra de tipo $0 - 1$, porque, $\mathbb{N}$ tiene un elemento distinguido cero y un endomorfismo $S : \mathbb{N} \rightarrow \mathbb{N}$. Este resultado se relaciona con el primer capítulo porque se consideró a $\mathbb{N}$ como una estructura algebraica.

Como $\mathbb{N}$ es un objeto de un álgebra de tipo $0 - 1$, entonces se procederá a mostrar que la minimalidad de $\mathbb{N}$ permite considerarlo como un objeto inicial. Después de presentar a los números naturales como un objeto inicial se procederá a escribir a los números reales como un objeto final en una categoría de coálgebras, esto permite evidenciar la dualidad entre $\mathbb{N}$ y $\mathbb{R}$.

3.3. $\mathbb{N}$ como Objeto Inicial en Teoría de Categorías

Se define a $\mathbb{N}$ como un álgebra de tipo $0 - 1$:

$$\mathbb{N} = \langle \mathbb{N}, s, 0 \rangle$$

- I. El objeto es el conjunto de números naturales.
- II. La operación es $s = \text{sucesor}^2$.
- III. El elemento distinguido es cero.
- IV. Un morfismo $s : \mathbb{N} \rightarrow \mathbb{N}$.

En el lenguaje de categorías se tiene a $\mathbb{N}$ como un álgebra $1 - 0$, porque:

- Se tiene un morfismo unario.

$$\mathbb{N} \xrightarrow{s} \mathbb{N}$$

- Se tiene un morfismo cero-ario.

$$f : A^0 \longrightarrow A$$

²S se denota para referirse al la función sucesor y s se denota para referirnos al morfismo, sin embargo, en esencia son lo mismo

$()$ representa al vacío.

$$A^0 = 1 = \{\emptyset\}$$

$$() \longrightarrow \mathbb{N}$$

$$1 \xrightarrow{0} \mathbb{N}$$

■ Combinando los morfismos anteriores se tiene:

$$\mathbb{N} + 1 \xrightarrow{s+0} \mathbb{N}$$

Por los ítem 1 y 2 se considera a los números naturales como un conjunto $\mathbb{N}$ dotado de una estructura, es decir, el conjunto $\mathbb{N}$ es un objeto de una álgebra de tipo $A = \langle A, \oplus, a^0 \rangle$. Normalmente en los textos se denomina este tipo de álgebras como $Alg_{1,0}$

Como $\mathbb{N}$ es un álgebra tipo $0 - 1$ $A = \langle A, \oplus, a^0 \rangle$ entonces se tiene:

$$A + 1 \xrightarrow{f+a} A$$

Los resultados en el ítem tres y el anterior se pueden resumir en el siguiente diagrama:

$$\begin{array}{ccc} \mathbb{N} + 1 & & A + 1 \\ \downarrow s+0 & & \downarrow f+a \\ \mathbb{N} & & A \end{array}$$

Por la definición de objeto inicial se plantea que existe un morfismo $h : \mathbb{N} \rightarrow A$ se puede encontrar otro morfismo $\mathbb{N} + 1 \rightarrow A + 1$, y es el morfismo $h + id$.

($id = identidad$)

$$\begin{array}{ccc}
\mathbb{N} + 1 & \xrightarrow{h+id} & A + 1 \\
\downarrow s+0 & & \downarrow f+a \\
\mathbb{N} & \xrightarrow{h} & A
\end{array}$$

El diagrama anterior conmuta porque si $n \in \mathbb{N} + 1$ entonces se cumple la siguiente igualdad $f + a(h + id(n)) = h(s + 0(n))$, esto implica que los números naturales se definen como un objeto inicial de un cierto tipo de categorías de álgebra $0 - 1$.

Se concluye que los números naturales se pueden construir inductivamente, lo cual según Dedkind en [6] permite establecer a $\mathbb{N}$ como un objeto inicial, además los números naturales son álgebras tipo $0 - 1$ porque tienen como operación unaria al sucesor y como operación cero-aria al elemento distinguido cero.

Observación 3.1 *Por ser $\mathbb{N}$ un objeto inicial cumple con dos principios:*

- *Principio de inducción = Método de demostración de propiedades.*
- *Principio de recursión = Garantiza la buena definición de operaciones.*

En el capítulo 1 se concluyó que los números naturales forman una estructura algebraica y en esta sección se mostrará que los números naturales son un objeto inicial desde la teoría de categorías, ahora bien las álgebras de tipo uno cero las denominaremos álgebras iniciales, que son una forma abstracta de la inducción, es por esto que se presentará a los números naturales como un algebra inicial. Entonces, un álgebra es inicial si existe un morfismo único que la relaciona con un álgebra arbitraria.

Definición 3.4 *Un álgebra es inicial, si para por una álgebra arbitraria existe un único homomorfismo de álgebras.*

$$(\text{álgebra inicial}) \xrightarrow[\text{(homomorfismo)}]{(\text{único})} (\text{álgebra arbitraria})$$

La anterior definición abre paso para evidenciar a los números naturales como un álgebra inicial, entonces:

Sea el conjunto $\mathbb{N}$ de los números naturales con las funciones cero y sucesor; $0 : 1 \rightarrow \mathbb{N}$ y $S : \mathbb{N} \rightarrow \mathbb{N}$. Estas funciones se combinan en una única función $[0, S] : 1 + \mathbb{N} \rightarrow \mathbb{N}$ del funtor $T(X) = 1 + X$.

Se mostrará que $[0, S] : 1 + \mathbb{N} \rightarrow \mathbb{N}$ es un álgebra inicial: para esto es necesario un homomorfismo único que permita relacionar las álgebras arbitrarias con un álgebra inicial.

Supongamos que tenemos un conjunto arbitrario U que lleva una estructura en forma de T-álgebra $[u, h] : 1 + U \rightarrow U$.

Se tiene que definir un homomorfismo mediador:

$$\begin{aligned} f : \mathbb{N} &\rightarrow U \\ f(n) &= h^n(U) \\ f(0) &= u \text{ y } f(n+1) = h(f(n)) \end{aligned}$$

Se tiene los siguientes morfismos:

$$[0, s] : 1 + \mathbb{N} \rightarrow \mathbb{N} \text{ y } [u, h] : 1 + U \rightarrow U$$

Estas dos funciones expresan el siguiente diagrama:

$$\begin{array}{ccc} 1 + \mathbb{N} & \xrightarrow{id+f} & 1 + U \\ \downarrow [0, S] & & \downarrow [u, h] \\ \mathbb{N} & \xrightarrow{f} & U \end{array}$$

Se puede observar que el anterior diagrama conmuta, lo que implica que $\mathbb{N}$ es una álgebra inicial. Como los números naturales son un objeto y álgebra inicial, entonces se definirán los principios de inducción y recursión en teoría de categorías.

3.3.1. Principio De Inducción

Como $\mathbb{N}$ es un objeto inicial en las $Alg_{0,1}$, se puede definir la inducción de la siguiente manera:

$$\begin{array}{ccc} \mathbb{N} & \xrightarrow{h} & A \\ s \downarrow & & \downarrow f \\ \mathbb{N} & \xrightarrow{h} & A \end{array}$$

Se debe verificar si la igualdad se cumple $f(h(n)) = h(s(n))$, con cero y cualquier elemento n .

Verifiquemos si se cumple con el cero, tal que $0 \in \mathbb{N}$:

$$\begin{array}{ccc} 0 & \xrightarrow{h} & h(0) \\ s \downarrow & & \downarrow f \\ s(0) & \xrightarrow{h} & A \end{array}$$

$f(h(0)) = h(s(0))$ como esta igualdad se cumple, entonces el diagrama conmuta.

Verifiquemos que se cumple para n

$$\begin{array}{ccc} n & \xrightarrow{h} & h(n) \\ s \downarrow & & \downarrow f \\ s(n) & \xrightarrow{h} & A \end{array}$$

Sea $n \in \mathbb{N}$ se le aplica el morfismo h , entonces se tiene $h(n)$ y se le aplica el morfismo f , se obtiene $f(h(n)) \in A$.

Por otro lado, $n \in \mathbb{N}$, si se aplica el morfismo s o el sucesor, entonces se tiene $s(n)$ y si se aplica el morfismo h , se tiene $h(s(n)) \in A$.

Entonces $f(h(n)) = h(s(n))$.

Ejemplo: utilizar la inicialidad para definiciones inductivas

Supongamos que queremos definir por la inducción la función $f(n) = 2^{-n}$ partiendo de los números naturales, para llegar a los números racionales.

Se definen las ecuaciones:

$$\begin{aligned} f(0) &= 1 \text{ (si } n = 0 \text{ en } f(n+1)) \text{ se tienen } f(0+1) = 1 \\ f(n+1) &= \frac{1}{2}f(n) \end{aligned}$$

Además, se define la función $f : \mathbb{N} \rightarrow \mathbb{Q}$ por inicialidad, ahora se tiene que poner una estructura de álgebra $1 + \mathbb{Q} \rightarrow \mathbb{Q}$ sobre el conjunto de los números racionales $\mathbb{Q}$.

La estructura se obtiene de l y s :

$$\begin{array}{ll} l : 1 \xrightarrow{1} \mathbb{Q} & s : \mathbb{Q} \xrightarrow{\frac{1}{2}(-)} \mathbb{Q} \\ * \mapsto \mathbb{Q} & x \mapsto \frac{1}{2}x \end{array}$$

Combinando l y s se tiene:

$$1 + \mathbb{Q} \xrightarrow{[1, \frac{1}{2}(-)]} \mathbb{Q} \text{ formando un álgebra de } \mathbb{Q}.$$

La función $f(n) = 2^{-n}$ se determina por inicialidad como una única función, haciendo que el diagrama conmute.

$$\begin{array}{ccc} 1 + \mathbb{N} & \xrightarrow{id+f} & 1 + \mathbb{Q} \\ \downarrow [0, s] & & \downarrow [1, \frac{1}{2}(-)] \\ \mathbb{N} & \xrightarrow{f} & \mathbb{Q} \end{array}$$

3.3.2. El principio de recursión.

[17] define a la recursión de la siguiente forma:

Definición 3.5 *Dados $f : 1 \rightarrow A$ y $g : \mathbb{N} \times \mathbb{N} \rightarrow A$, entonces existe una única función $h : \mathbb{N} \rightarrow A$ tal que:*

$$I. h(0) = f(0)$$

$$II. h(n+1) = g(n, h(n))$$

Como la recursión permite establecer que la operación suma este bien definida, entonces combinando la inducción con la recursión se tiene que:

$$+ : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$$

$$\text{tal que: } +(n, m) = +_n(m)$$

Si se toma a $f : 1 \rightarrow A$ y se opera con el morfismo $g : \mathbb{N} \times \mathbb{N} \rightarrow A$, se obtiene:

$$1 + \mathbb{N} \times \mathbb{N} \xrightarrow{[f, g]} A$$

Como se obtuvo que $\mathbb{N}$ es un objeto inicial, se obtiene:

$$1 + \mathbb{N} \xrightarrow{[s, 0]} \mathbb{N}$$

Empleando el teorema de recursión en un diagrama de teoría de categorías se tiene:

$$\begin{array}{ccc}
1 + \mathbb{N} & \xrightarrow{[s,0]} & \mathbb{N} \\
\downarrow [id_1, \langle id_{\mathbb{N}}, h \rangle] & & \downarrow h \\
1 + \mathbb{N} \times \mathbb{N} & \xrightarrow{[f,g]} & A
\end{array}$$

En el diagrama anterior se observa que se cumplen los dos puntos de la definición:

- I. Para $0 \in 1 + \mathbb{N}$, entonces: $h \circ [0, s](0) = h(0) = [f, g] \circ [id_{1\circ}, \langle id_{\mathbb{N}}, h \rangle](0) = f(0)$.
- II. Para $n \in 1 + \mathbb{N} (n \neq 0)$, se tiene que: $h \circ [0, s](n) = h \circ s(n) = h(n+1) = [f, g] \circ [id_{1\circ}, \langle id_{\mathbb{N}}, h \rangle](n) = [f, g] \circ [id_{1\circ}, \langle n, h(n) \rangle] = g(n, h(n))$.

Hasta al momento se ha caracterizado a los números naturales como un objeto y un álgebra inicial desde una visión categórica con los principios de inducción y recursión. Para cerrar esta sección se retoma que Dedekind define a los números naturales como el mínimo conjunto inductivo, lo cual implica que dicha minimalidad significaba que $\mathbb{N}$ es un objeto inicial, y por ende cumple con los principios de inducción y recursión. En lo que sigue del capítulo dualmente se presentarán a los números reales como un objeto final, que cumplen con los principios de coinducción y corecursión. Para cumplir este fin se presentará a $\mathbb{R}$ como una coálgebra.

3.4. Coálgebra y Coinducción

Hilbert en 1900 caracterizó a los números reales como el máximo cuerpo arquimediano totalmente ordenado, desde este momento se puede asociar a los números reales como un objeto final [17]. Por otro lado, el trabajo de Peter Aczel de hiperconjuntos y el axioma de anti-fundamentación abrieron paso a los conceptos trabajados en esta sección del capítulo, como las coálgebras, que son duales a las álgebras, y de igual forma con los principios de coinducción y corecursión, que son duales a los principios de inducción y recursión. Con los anteriores hechos históricos se promueve evidenciar la dualidad de $\mathbb{N}$ y $\mathbb{R}$ mediante el lenguaje de categorías, es por esto que en este capítulo se

hará un estudio de los números reales como objetos y coálgebras finales.

A continuación se definirá las coálgebras, con el fin de reconocer a $\mathbb{R}$ como un objeto final.

Observación 3.2 *Dado un funtor fijo T , se forma una categoría de todas la T -coálgebras.*

Definición 3.6 *Para un funtor T , una coálgebra (o una T -coálgebras) es un par (U, c) que consiste en un conjunto U y una función $c : U \rightarrow T(U)$. Al igual que para las álgebras, tiene un conjunto U que es el dominio y una función estructural c permitiendo el funcionamiento de la coálgebra (U, c) .*

Entonces, ¿cuál es la diferencia entre un álgebra $T(U) \rightarrow U$ y una coálgebra $U \rightarrow T(U)$? En esencia es la diferencia que existe entre la construcción y la observación.

Un *álgebra* consiste en un módulo que configura a un conjunto U con una función $T(U) \rightarrow U$ (va en este U), esto implica que se pueden construir los elementos de un conjunto teniendo una función o un operador y un conjunto base o de dominio, por ejemplo: se tiene como conjunto base, el de los números naturales y el operador sucesor $S(n)$, si se aplica $S(\mathbb{N}) \rightarrow \mathbb{N}$ se tiene al mismo $\mathbb{N}$

Una *coálgebra* consta de un conjunto U de partida con una función $U \rightarrow T(U)$ en la dirección opuesta de las álgebras, saliendo del conjunto U para llegar a un conjunto $T(U)$. Aquí no se construye, si no que se analiza los elementos del conjunto de llegada $T(U)$. En este caso no se sabe cómo formar los elementos en U , sólo se conocen las operaciones que actúa en U , lo que permite obtener información acerca de U , lo que implica que se tiene un acceso limitado a U .

Definición 3.7 *Dados $V \subseteq X$, e $i : V \rightarrow X$ tal que el operador i es un morfismo, entonces V es una subcoálgebra de X .*

$$\begin{array}{ccc}
 V & \xrightarrow{i} & X \\
 \gamma_V \downarrow & & \downarrow \gamma_X \\
 \mathbf{F}V & \xrightarrow{\mathbf{F}i} & \mathbf{F}X
 \end{array}$$

Donde se cumple $\mathbf{F}i \circ \gamma_V = \gamma_X \circ i$

Ejemplo

Sean

$$A = \{1, 2, 3, 4, 5, \dots, n\}$$

$$B = \left\{ \frac{1}{2}, \frac{1}{4}, \frac{1}{6}, \frac{1}{8}, \frac{1}{10}, \dots, ? \right\}$$

Los conjunto A y B están relacionados mediante $\{1 \text{ con } \frac{1}{2}\}, \{2 \text{ con } \frac{1}{4}\}$ y así sucesivamente, si se pudiera hallar la función con la cual se logra obtener B , lo que se debe hacer es tomar cada elemento del conjunto B y analizarlo para poder comprender la cola infinita, de esta forma, lo que se está haciendo es tratar de crear una función la cual obtenga a B , a este procedimiento se le conoce como coinducción, dado que se hace análisis de los objetos. Retomando el ejemplo, es evidente que la función que crea a B partiendo desde A es $\frac{1}{2n}$, sin embargo esta función u operador se halló teniendo en cuenta el inicio del conjunto y su parte infinita o su cola.

Como se había mencionado las álgebras implementan operaciones constructoras que son la misma definición de recursividad, en cambio las coálgebras implementan las operaciones de observación o de transmisión de funciones; estas observan elementos de datos o la estructura de un conjunto.

Ejemplo

Sea el funtor $T(X) = A \times X$, donde A es un conjunto fijo. La coálgebra $T \rightarrow T(U)$ se compone de dos funciones $U \rightarrow A$ y $U \rightarrow U$.

Obsérvese los siguientes diagramas:

$$U \xrightarrow{\text{next}} U$$

$$U \xrightarrow{\text{value}} A$$

$$\begin{array}{ccc} U & \xrightarrow{\text{next}} & A \times U \\ & \nearrow \text{value} & \\ U & & \end{array}$$

Llamados antes $\text{value} : U \rightarrow A$ y $\text{next} : U \rightarrow U$. Entonces a un elemento $u \in U$ se le puede realizar las operaciones:

- (1). producir un elemento en A , es decir, $\text{value}(u)$
- (2). producir un siguiente elemento en U , que es igual a $\text{next}(u)$.

Ahora se puede repetir (1) y (2) y formar un elemento de A , siendo la composición de $\text{value}(\text{next}(u))$.

Al tomar $u \in U$, y se realiza la operación en (1) $\text{value}(u)$, se obtiene un $a_1 \in A$. Ahora se realiza la segunda (2) o la operación next en u , es decir $\text{next}(u)$, y se produce un elemento u_2 que pertenece a U . Luego se opera value en u_2 obteniéndose a_2 , que pertenece a A , entonces los a_n elementos de A , tiene la siguiente forma $a_n = \text{value}(\text{next}^n(u))$, mediante este procedimiento se puede obtener cada elemento $u \in U$, siendo una secuencia infinita $(a_1, a_2) \in A^{\mathbb{N}}$ de los elementos de A . La secuencia de elementos que U origina, es lo que se puede observar. Dos elementos $(u_1, u_2) \in U$ también puede dar lugar a la misma secuencia de elementos de A .

Por tanto, los elementos de A se pueden estudiar a través de los elementos de U , dado que del conjunto U origina al conjunto A y esto hace que los elementos de A y los elementos de U se comporten de la misma forma y

cumplan las mismas condiciones.

Con el anterior ejemplo se puede entender que los funtores son esencialmente morfismos que permiten comparar coálgebras. Esto implica que se puede definir un morfismo entre coálgebras.

Definición 3.8 Sea T un funtor y c_1 y c_2 coálgebras.

Un homomorfismo de coálgebras (o mapa de coálgebras) entre una T -coálgebra $U_1 \xrightarrow{c_1} T(U_1)$ y otra $U_2 \xrightarrow{c_2} T(U_2)$ es una función $f : U_1 \rightarrow U_2$ que conmuta en las operaciones $c_2 \circ f = T(f) \circ c_1$.

$$\begin{array}{ccc} U_1 & \xrightarrow{f} & U_2 \\ c_1 \downarrow & & \downarrow c_2 \\ T(U_1) & \xrightarrow{T(f)} & T(U_2) \end{array}$$

Definición 3.9 Decimos que una coálgebra $\langle Z, d \rangle$ es final, si para cualquier otra F -coálgebra $\langle U, c \rangle$ existe un único morfismo $f : U \rightarrow Z$

Una coálgebra final $d : Z \rightarrow T(Z)$ es una coálgebra tal que para cada coálgebra $c : U \rightarrow T(U)$ es un único mapa de coálgebras $(U, c) \rightarrow (Z, d)$.

$$\begin{array}{ccc} U & \xrightarrow{f} & Z \\ c \downarrow & & \downarrow d \\ T(U) & \xrightarrow{T(f)} & T(Z) \end{array}$$

Una coálgebra es final o terminal si por una coálgebra arbitraria existe un único homomorfismo de coálgebras:

$$(coálgebras\ arbitrarias) \xrightarrow[\text{(homomorfismo)}]{\text{(único)}} (coálgebras\ finales)$$

Con las dos definiciones anteriores se evidencia que se puede relacionar una coálgebra arbitraria con una coálgebra final por medio de la existencia de un único morfismo. Este es un paso importante para poder clasificar a los números reales como coálgebras finales.

Ejemplo de coálgebra final

Para un conjunto fijo A , se considera el funtor $T(X) = A \times X$. Se afirma que la coálgebra final de este funtor es el conjunto $A^{\mathbb{N}}$ de listas infinitas de elementos de A , con una estructura coalgebraica :

$$\langle head, tail \rangle : A^{\mathbb{N}} \rightarrow A \times A^{\mathbb{N}}$$

dada por

$$head(\alpha) = \alpha(0) \text{ y } tail(\alpha) = \lambda x. \alpha(x + 1)$$

Por lo tanto $head$ toma el primer elemento de una secuencia infinita $(\alpha(0), \alpha(1), \alpha(2) \dots)$ de elementos de A , y $tail$ toma la lista restante. Se evidencia de que el par de funciones $\langle head, tail \rangle : A^{\mathbb{N}} \rightarrow A \times A^{\mathbb{N}}$ es un isomorfismo. Se afirma que para una coálgebra arbitraria $\langle head, tail \rangle : U \rightarrow A \times U$, existe un único morfismo de coálgebra $f : U \rightarrow A^{\mathbb{N}}$; y está dado para $u \in U$ y $n \in \mathbb{N}$ por:

$$f(u)(n) = value(next^{(n)}(u))$$

De hecho $head \circ f = value$ y $tail \circ f = f \circ next$, hacen a f un morfismo de coálgebras y como f es único, entonces satisface las dos igualdades. Esto se puede comprobar fácilmente:

Como $u \in U$ es una lista infinita de elementos de A que surge como $value(u), value(next(u)), value(next^{(n)}(u)), \dots$, Ahora vemos que este comportamiento observable de u es precisamente el resultado $f(u) \in A^{\mathbb{N}}$ en u , con un único morfismo f para la coálgebra final. Por lo tanto los elementos de la coálgebra final dan el comportamiento observable. Esto es un ejemplo típico de coálgebras finales.

De el ejemplo anterior se puede concluir que una coálgebra es final si existe un único morfismo que va de una coálgebra arbitraria a la coálgebra final. Además, una coálgebra tiene dos operadores los cuales se denominan $head$

y *tail*, y permiten observar una lista infinita de datos. El primer operador *head* permite observar la cabeza de una lista infinita y *tail* observa el resto de la lista.

Ahora bien se puede plantear un diagrama del ejemplo anterior:

$$\begin{array}{ccc}
 U & \xrightarrow{f} & A^{\mathbb{N}} \\
 \downarrow \langle value, next \rangle & & \downarrow \langle head, tail \rangle \\
 A \times U & \xrightarrow{\langle A \times f \rangle} & A \times A^{\mathbb{N}}
 \end{array}$$

Si se cumple la siguiente igualdad:

$(\langle A \times f \rangle \circ \langle value, next \rangle)(u) = (\langle head, tail \rangle \circ f)(u)$ significa que el diagrama conmuta, entonces se dice que la coálgebra es final.

3.4.1. $\mathbb{R}$ como Objeto Final y Coálgebra Final

Con el fin de ver a $\mathbb{R}$ como un objeto y una coálgebra final, se presentará un ejemplo tomado de [23], el cual introduce a los números reales pertenecientes al intervalo $[0, 1)$, y luego se revisarán los conceptos involucrados en tal ejemplo.

Ejemplo

Consideremos una máquina con un conjunto de estados S y una aplicación que permite las transiciones o el paso de un estado a otro, llamada c . Ahora, supongamos que luego de hacer la acción c , la máquina tiene todas o alguna de las siguientes opciones: mostrar un comportamiento (que podría ser un sonido, una letra impresa, etc.) o brindar la posibilidad de una modificación (como el teclado de un computador). Estas posibilidades de observación y modificación es lo que en computación se conoce como *entrada y salida* o por sus siglas en inglés *I/O*. Para admitir ese tipo de interacción, lo adecuado sería cambiar el espacio de estados de llegada luego de aplicar c . Este espacio con nuevas posibilidades de interacción lo vamos a llamar $F(S)$.

Supongamos que la máquina anterior tenga como espacio de estados los números reales pertenecientes al intervalo $[0, 1)$, y se necesita que muestre su

representación decimal, además, de considerar si su representación es finita o infinita. De esta manera las posibles observaciones en cada interacción pertenecen al conjunto $A = \{0, 1, 2, 3, 4, 5, 6, 7, 9\}$. Sea r el número real en cuestión, así la coálgebra actúa de la siguiente manera:

$$c(r) = \begin{cases} \perp & \text{si } r = 0 \\ (d, 10r - d) & \text{en otro caso, donde } d \text{ es tal que } d \leq 10r \leq d + 1 \end{cases} \quad (3.1)$$

La pareja ordenada $(d, 10r - d)$ tiene como significado calcular la cabeza y la cola del decimal entre $[0, 1)$

El símbolo $\perp$ significa que la sucesión es finita y por lo tanto debe parar, como este resultado no pertenece ni a A ni a $[0, 1)$ se puede organizar como un elemento independiente de los anteriores, esto se puede hacer mediante el coproducto, finalmente el espacio de llegada es: $F([0, 1)) = A \times [0, 1) \cup \perp$.

La manera como se expresa su resultado (d, r_2) quiere decir que el espacio después de aplicar la transición es el producto cartesiano $A \times [0, 1)$.

Se aplicará $\frac{1}{4}$ en la siguiente proceso c :

$$c(r) = (d, 10r - d)$$

$$r = \frac{a}{b}$$

Para realizar el proceso es necesario conocer el valor de d el cual se halla de la siguiente manera:

$$\begin{array}{r} 1 \mid 4 \\ 2 \quad 0,25 \end{array}$$

$$\frac{1}{4} = \boxed{0,2} 5$$

Se estudia en la primera transición de estado a 0,2 dejando a un lado a la cola que es 5 para la próxima transición c :

$0,2 \times 10 = 2$ obteniendo $d = 2$ en la primer preimagen de la pareja ordenada.

Ahora es necesario hallar la imagen de la primera pareja ordenada y para ello se aplica el proceso c :

$$c(\frac{1}{4}) = (2, \frac{10}{4} - 2)$$

$$c(\frac{1}{4}) = (2, \frac{1}{2}) \text{ ya que } 2 \leq \frac{5}{2} < 3$$

En el primer estado $\frac{1}{4} = 0,25$, al hacer la transición se estudió de la parte decimal 0,2 que es la cabeza del decimal 0,25, ahora se estudiará la parte decimal 0,5 que es la cola; es decir, que en cada estudio o aplicación se toma la parte entera y el primer decimal.

Entonces la aplicación c se repite nuevamente con $\frac{1}{2}$

si repetimos el proceso con $\frac{1}{2}$ tenemos:

$$c(\frac{1}{2}) = (5, 0) \text{ ya que } 5 \leq 5 < 6$$

se aplica el proceso en 0 y este cumple con la primera cláusula de la coálgebra terminando la aplicación c .

$$c(0) = \perp$$

De esta manera, las repetidas aplicaciones de c a un determinado número en $[0, 1)$ muestra su representación decimal paso a paso, en este caso muestra 0,25 y luego detiene el cómputo, representando a $\frac{1}{4}$.

El siguiente es un ejemplo de un número con representación infinita:

Se utilizará $\frac{1}{3}$ en proceso c .

Primero se hallará d .

$$\begin{array}{r|l} 1 & 3 \\ 1 & 0,3 \end{array}$$

$$\frac{1}{3} = \boxed{0.3} 33\ldots \quad \frac{1}{3} = 0. \bar{3}$$

$0,3 \times 10 = 3$ obteniendo $d = 3$ en la primer preimagen de la pareja ordenada, el cual nos permite calcular la cabeza del decimal $0. \bar{3}$.

Para hallar la imagen es necesario aplicar la definición $(d, 10r - d)$

$$(10 \cdot \frac{1}{3}) - 3 = \frac{1}{3}$$

$$c(\frac{1}{3}) = (3, \frac{1}{3}) \text{ ya que } 3 \leq \frac{10}{3} < 4$$

Se realiza la transición c nuevamente con $c(\frac{1}{3})$, puesto que este decimal permite calcular la cola, obteniendo:

$$c(\frac{1}{3}) = (3, \frac{1}{3}) \text{ ya que } 3 \leq \frac{10}{3} < 4$$

Como se puede observar obtenemos nuevamente la transición $c(\frac{1}{3})$, puesto que el primer dígito después de la coma es 3 y el siguiente es 3, y así periódicamente; entonces se tiene un proceso c infinito que calcula una cola infinita.

¿ Donde quedan los irracionales con este ejemplo?

Supóngase que $r = \frac{1}{\pi}$ y se quiere representar en su forma decimal mediante la transición c , entonces para estudiarlo es necesario utilizar $(d, 10r - d)$ donde $d = 3$ y $\frac{10}{\pi} - 3 = 0,183098\ldots$, entonces se tiene la pareja $(3, \frac{10-3\pi}{\pi})$, si se continua el proceso con $r = \frac{10-3\pi}{\pi}$ se obtiene la pareja $(1, \frac{10^2-31\pi}{\pi})$ el proceso continua estudiando las cabezas y colas infinitas de $\frac{1}{\pi}$, arrojando cabezas pertenecientes al conjunto A y colas pertenecientes a los irracionales que se encuentran en el intervalo $[0, 1)$.

Este ejemplo muestra la versatilidad de un cambio de espacio de estados, por un lado se pueden hacer observaciones de lo que está ocurriendo y por otro lado se tiene la posibilidad de frenar el proceso si las condiciones lo requieren. También está la opción de representar un programa que corra por siempre, como por ejemplo un sistema operativo.

El ejemplo anterior se puede identificar como una coálgebra, dado que: la máquina tiene como espacio de estados a los números reales pertenecientes al intervalo $[0, 1)$, y la tarea era expresar la representación decimal de cada número perteneciente a este intervalo, por medio de la función $f([0, 1)) = A \times [0, 1) \cup \perp$

Para expresar este ejemplo como una coálgebra primero se tienen que definir los dos posibles resultados de la máquina:

- El primero es cuando el proceso finaliza: Este proceso permite obtener un conjunto A^* de sucesiones finitas de los elementos de A .
- El segundo resultado es cuando el programa de la máquina corre por siempre: Este proceso permite obtener un conjunto $A^{\mathbb{N}}$ de sucesiones infinitas de los elementos de A .

Los conjuntos A^* y $A^{\mathbb{N}}$ permiten definir el conjunto A^ω siendo $\omega = \{0, 1, 2, \dots\}$ de la siguiente manera:

$$A^\omega = A^* \cup A^{\mathbb{N}}$$

A este conjunto se le asignará una estructura de transición llamada $\langle head, tail \rangle$ formando la siguiente coálgebra de la categoría $CoAlg(F)$:

$$\begin{array}{c} A^\omega \\ \downarrow \langle head, tail \rangle \\ \{\perp\} \cup (A \times A^\omega) \end{array}$$

Para comprender como la coálgebra funciona en el ejemplo, se observa el siguiente caso particular en el que $r = \frac{1}{4}$ obteniendo:

$$c(\frac{1}{4}) = (2, \frac{1}{2})$$

Se puede observar que $\frac{1}{4} \in A^\omega$ y al aplicarle la transición c se obtiene la pareja $(2, \frac{1}{2})$ en la cual $2 \in \{\perp\} \cup A$ y $\frac{1}{2} \in A^\omega$

En esta coálgebra se tiene que $\langle head, tail \rangle$ es la transición de un estado a otro que contiene los operadores de observación h para estudiar la cabeza y t para estudiar la cola de cualquier decimal. Formando así la siguiente coálgebra.

$$\begin{array}{c} A^\omega \\ \downarrow \langle h, t \rangle \\ \{\perp\} \cup A \times A^\omega \end{array}$$

Esta coálgebra contiene dos posibles resultados de la máquina de la siguiente forma:

$$\sigma \mapsto \begin{cases} \perp & \text{si } \sigma \text{ es la sucesión vacía } \langle \rangle \\ (a, \sigma') & \text{si } \sigma = a \cdot \sigma' \text{ con cabeza } a \in A \text{ y cola } \sigma' \in A^\omega \end{cases} \quad (3.2)$$

Observación 3.3 *La aplicación $\langle h, t \rangle$ envía $\perp$ a la sucesión vacía y corresponde al proceso de la máquina cuando finaliza. Por otro lado, esta aplicación envía a $(a, \sigma) \in A \times A^\omega$ a la sucesión a, σ que corresponde al proceso que corre por siempre de la máquina.*

Ahora es necesario probar que $\langle h, t \rangle : A^\omega \rightarrow \{\perp\} \cup (A \times A^\omega)$ es una coálgebra final, es decir:

Que para cada coálgebra arbitraria $\langle k, s \rangle : S \rightarrow \{\perp\} \cup (A \times S)$ sobre un conjunto S existe un único homomorfismo $[\langle k, s \rangle] : S \rightarrow A^\omega$ entre las coálgebras. Para verificar lo anterior es necesario comprobar que el siguiente diagrama conmuta.

$$\begin{array}{ccc}
S & \xrightarrow{[\langle k, s \rangle]} & A^\omega \\
\downarrow \langle k, s \rangle & & \downarrow \langle h, t \rangle \\
\{\perp\} \cup A \times S & \xrightarrow{\{\perp\} \cup A \times [\langle k, s \rangle]} & \{\perp\} \cup A \times A^\omega
\end{array}$$

Para determinar si el diagrama conmuta:

Sea $n \in S$ entonces, se verifica la siguiente igualdad:

$$(\{\perp\} \cup A \times [\langle k, s \rangle] \circ \langle k, s \rangle)(n) = (\langle h, t \rangle \circ [\langle k, s \rangle])(n)$$

- I. Como $n \in S$ y al aplicarle el operador $\langle k, s \rangle$ se obtiene la pareja (a, n) , en la cual $a \in \{\perp\} \cup A$ y $(a, n) \in \{\perp\} \cup A \times S$. Por último, si se aplica el operador $\{\perp\} \cup A \times [\langle k, s \rangle]$ a (a, n) , se tiene que $(a, n) \in \{\perp\} \cup A \times A^\omega$
- II. Por otro lado, sea $n \in S$ y se le aplica el operador $[\langle k, s \rangle]$ a n y como el operador es el morfismo establecido anteriormente se tiene que $n \in A^\omega$. Ahora si a n se le aplica el operador $\langle h, t \rangle$, se obtiene: $(a, n) \in \{\perp\} \cup A \times A^\omega$.

De esta forma se puede concluir que el diagrama conmuta, lo cual implica que $\langle h, t \rangle : A^\omega \rightarrow \{\perp\} \cup (A \times A^\omega)$ es una cóalgebra final para las de su tipo.

El ejemplo anterior, muestra que se puede representar en su forma decimal los números reales pertenecientes en el intervalo $[0, 1)$, y así son una cóalgebra final, por lo que se puede concluir que el conjunto de los números reales son un objeto final como lo mencionan Ortíz y Valencia en [17], puesto que Hilbert en 1900 definió que eran un cuerpo maximal arquimediano totalmente ordenado. Teniendo en cuenta el resultado anterior se puede representar a $\mathbb{R}$ en el siguiente diagrama.

$\mathbb{R}$ como Coálgebra Final

$$\begin{array}{ccc}
S & \xrightarrow{[\langle k, s \rangle]} & \mathbb{R}^\omega \\
\downarrow \langle k, s \rangle & & \downarrow \langle h, t \rangle \\
\{\perp\} \cup A \times S & \xrightarrow{\{\perp\} \cup A \times [\langle k, s \rangle]} & \{\perp\} \cup A \times \mathbb{R}^\omega
\end{array}$$

Con el diagrama anterior se puede observar que los números reales en intervalo $[0, 1)$, se pueden representar en forma decimal, dado que existe un único morfismo $[\langle k, s \rangle]$ que relaciona las coálgebras.

Retomando el capítulo 2, los números naturales se consideraron como el mínimo conjunto inductivo y en este capítulo tres se concluyó que dicha minimalidad define a $\mathbb{N}$ como un objeto inicial, dualmente, Hilbert en 1900 clasificó a los números reales como el máximo cuerpo arquimediano totalmente ordenado, dicha maximalidad define a los números reales como un objeto final. Por otro lado, si los naturales son un conjunto inductivo entonces dualmente los números reales son un conjunto coinductivo que cumple la siguiente definición:

$$\nu F = \bigcup \{X \subseteq L : X \subseteq F(X)\}$$

Con el fin de evidenciar la dualidad entre $\mathbb{R}$ y $\mathbb{N}$ se presenta la ecuación de Pavlović.

3.5. La Ecuación de Pavlović

En los capítulos anteriores, se ha estudiado a los números naturales como un objeto inicial y a los números reales como un objeto final; es decir, se ha intentado estudiar a $\mathbb{R}$ mediante la dualidad que tiene con $\mathbb{N}$.

Cuando se establece a $\mathbb{N}$ como un objeto inicial se define el principio de inducción y de recursión, dualizando lo anterior si se establece a $\mathbb{R}$ como un

objeto final se evidencia el principio de coinducción y corecursión. Con lo anterior [19] presenta la siguiente ecuación:

$$\frac{\text{Inducción}}{\text{Aritmética}} \approx \frac{\text{Coinducción}}{\text{Análisis}}$$

Como se puede observar en el lado izquierdo, las infinitas construcciones en la aritmética elemental tienen como base el principio de inducción, un ejemplo de ello son los números naturales que se definen recursivamente con un elemento distinguido que es el cero (vacío), y si n pertenece a $\mathbb{N}$, el sucesor de n también. La pregunta sería ¿Para qué se utiliza la aritmética? y la respuesta se halla en la función sucesor, puesto que tal función parte de un elemento n y para obtener el conjunto $\mathbb{N}$ se necesita de $f : n + 1$ y allí se está utilizando la aritmética. Cuando el objeto es inicial entonces se hace uso de la inducción y además se hace necesario establecer el principio de recursión que permite establecer que las operaciones están bien definidas, es por ello que esta parte de la ecuación quedaría así:

$$\frac{\text{Inducción} + \text{Recursión}}{\text{Aritmética}}$$

Dualmente al observar el lado derecho, se tiene que el cálculo elemental hace uso de la coinducción, puesto que al momento de estudiar un objeto final es necesario conocer la estructura que lo compone y esto lo permite el principio de coinducción, puesto que hace uso de operadores de observación, además también es importante el principio de corecursión porque permite establecer que las operaciones estén bien definidas.

En este capítulo se presentó un ejemplo que evidencia lo anterior, el cual planteaba a una máquina con un conjunto de estados S y una aplicación c y tenía como propósito representar a los números en el intervalo $[0, 1)$ en su forma decimal. Al aplicar la transición c se usaba a los operadores de observación *head* y *tail* para pasar de un estado al otro, esto implicaba el uso del análisis. Para evidenciar lo anterior veamos un caso particular del ejemplo: Cuando $r = \frac{1}{4}$ y se obtenía la pareja $(2, \frac{1}{2})$ y se realizaba otra vez la transición c con $r = \frac{1}{2}$ y se obtuvo $(5, 0)$ allí se está utilizando el análisis,

porque se estudia la cabeza y la cola de $r = \frac{1}{4}$.

Con los resultados anteriores se reformula la ecuación de Pavlović y Pratt en [19] de la siguiente forma:

$$\frac{\text{Inducción} + \text{Recursión}}{\text{Aritmética}} \approx \frac{\text{Coinducción} + \text{Corecursión}}{\text{Análisis}}$$

La ecuación de Pavlović reformulada abre paso al estudio de algunos contenidos del cálculo que tienen inmerso el razonamiento coinductivo.

3.6. La Coinducción en el Cálculo

En la enseñanza de los números reales es necesario conocer el método o principio que permite el estudio de este objeto final, es por ello que en este trabajo de grado se ha propuesto otra forma de ver a los números reales, como lo menciona [18] cuando se enseña de una manera adecuada los temas relacionados con los números reales, como el cálculo, se puede observar que se desarrolla un razonamiento coinductivo. Es por ello que de aquí en adelante se observará el razonamiento coinductivo que contiene las nociones matemáticas que estudian los números reales, como la integral y las series de Taylor, con el fin de obtener otras herramientas para la enseñanza de este objeto matemático.

Según [18] los elementos del cálculo poseen razonamiento coinductivo, y para empezar a identificarlo el punto de partida es la observación de las estructuras algebraicas de las listas infinitas (Streams³) y para ello se necesita definir las siguientes operaciones básicas, dada una función:

$$\text{head}(f) = f(0)$$

$$\text{tail}(f) = f'$$

³Las Streams de acuerdo con [7] son listas infinitas de datos, este término se empleará en lo que sigue del documento.

$$a :: f = (x \mapsto a + \int_0^x f)$$

Nota: el simbolo $::$ se usa para separar la cabeza y la cola de una sucesión.

Aquí , *head* y *tail* son operaciones de observación de una coálgebra de listas infinitas.

Las propiedades básicas que contienen las *streams* de álgebras que capturan gran parte del cálculo son:

$$head(a :: \beta) = a \quad (1)$$

$$tail(a :: \beta) = \beta \quad (2)$$

$$head(a) :: tail(\beta) = \alpha \quad (3)$$

Las anteriores propiedades se pueden ver de la siguiente forma:

Cuando se analiza (1) $head(a :: \beta)$ lo que realmente se está realizando es; la aplicación de las reglas de la integral definida con respecto a un intervalo.

$$\begin{aligned} head(a :: f) &= head(a + \int_0^x f(t)dt) \\ &= (a + \int_0^x f(t)dt)(0) \\ &= a + \int_0^0 f(t)dt = a \end{aligned}$$

Cuando se observa (2) $tail(a :: \beta)$ lo que se tiene es la aplicación de la integral con respecto a la función subintegral:

$$\begin{aligned} tail(a :: f) &= tail(a + \int_0^x f(t)dt) \\ &= \frac{d}{dx}(a + \int_0^x f(t)dt) \\ &= 0 + \frac{d}{dx} \int_0^x f(t)dt = f \end{aligned}$$

Por último cuando se analiza $(3)head(a) :: tail(\beta)$ se tiene el teorema fundamental del cálculo:

$$head(a) :: tail(f) = f(0) :: f' = f(0) + \int_0^x f'(t)dt = f$$

Ejemplo particular

Sea la $f(t) = \cos(t)$

Para calcular $(head)$ se define la operación y se realiza la propiedad anteriormente mencionadas:

- La operación:

$$head(f) = f(0)$$

$$f(0) = \cos(0) = 1$$

- La propiedad (1):

$$head(a :: \beta) = a + \int_0^0 f(t)dt = a$$

Como $x \rightarrow a$ entonces:

$$a + \int_0^0 \cos(t)dt = a + \left[\sin(t) \right]_0^0 = a + \sin(0) - \sin(0) = a$$

Por otro lado, cuando se analiza $tail$ en $f(t)$ se realiza lo siguiente:

- La operación:

$$\text{tail}(f) = f'$$

$$f'(t) = -\text{sen}(t)$$

- La propiedad (2):

$$\begin{aligned} \text{tail}(a :: f) &= 0 + \frac{d}{dx} \int_0^x f(t) dt = f \\ &= 0 + \frac{d}{dx} \int_0^x \cos(t) dt = 0 + \frac{d}{dx} \text{sen}(t) \Big|_0^x = 0 + \cos(x) = \cos(x) \Rightarrow f \end{aligned}$$

Por ultimo se observará la propiedad (3) que contiene el teorema fundamental del cálculo:

$$\text{head}(a) :: \text{tail}(\beta) :: f' = f(0) + \int_0^x f'(t) dt = f$$

$$\begin{aligned} &= 1 + \int_0^x -\text{sen}(t) = 1 + \cos(t) \Big|_0^x = 1 + \cos(x) - \cos(0) \\ &= 1 + \cos(x) - 1 = \cos(x) \Rightarrow f \end{aligned}$$

¿Que se obtiene si se aplica varias veces la propiedad (3)?

Si la propiedad (3) se repite varias veces; es decir, si se calcula la cabeza y la cola de la *stream* varias veces con la función f se obtiene:

$$f = f(0) :: f'(0) :: \dots :: f^n(0)$$

$$f = \underbrace{f(0)}_{\text{cabeza}} \underbrace{f'(0) :: \dots :: f^n(0)}_{\text{cola}}$$

Con esta sucesión se tiene la expresión de Taylor, pero como este procedimiento está centrado sobre el punto cero, se le denomina expansión de McLaurin.

$$f(x) = f(0) + f'(0)x + \dots + \frac{f^{(n)}(0)x^n}{n!} +$$

Cuando se aplica infinitas veces (3) se obtiene una *stream* que captura la noción de cólgebras, porque si se analizan las cabezas y las colas de la *stream*, se tiene un conjunto de secuencias infinitas de datos que forman una cólgebra final:

$$f = \underbrace{f(0)}_{\text{primer dato}} :: \underbrace{f'(0)}_{\text{segundo dato}} :: \dots :: \underbrace{f^{(n)}(0)}_{\text{dato infinito}}$$

Como la secuencia infinita es una Stream de un cólgebra, entonces la expansion de Taylor se obtiene utilizando un único homomorfismo:

$$(\text{cólgebras arbitrarias}) \xrightarrow[\text{(homomorfismo)}]{(\text{único})} (\text{cólgebras finales})$$

$$(\text{stream funciones analíticas}) \xrightarrow[\text{(homomorfismo)}]{(\text{único})} (\text{cólgebras finales})$$

$$(f(x)) \xrightarrow[\text{(homomorfismo)}]{(\text{único})} (\text{expansión de Taylor})$$

Ejemplo

Un ejemplo de lo anterior de una *stream* que captura la noción de cólgebra es el que presenta [18], el cual trata de un programa que emite una lista o stream de datos que forman una secuencia infinita:

Sea $f(x) = \text{sen}(x)$

$$f'(x) = \cos(x); \cos(0) = 1$$

$$f''(x) = -\text{sen}(x); -\text{sen}(0) = 0$$

$$\vdots$$

$$f'''(x) = \text{sen}(x) = f$$

Con los anteriores resultados se tienen:

$$f(0) = 0, f'(0) = 1, f''(0) = 0, f'''(0) = -1 \text{ y } f''' = f$$

Estos resultados forman la secuencia de datos de la *stream*:

$$f = 0 :: 1 :: -1 :: f$$

3.6.1. La Serie de Taylor desde La Coinducción

Brook Taylor en su trabajo “Methodus Incrementorum Directa et inversa” (1715) desarrolló lo que hoy se conoce como cálculo de las diferencias finitas. Este trabajo contenía la famosa fórmula conocida como el Teorema de Taylor, cuya importancia se reconoció en 1172, cuando Joseph Louis Lagrange lo definió como: “El fundamento principal del cálculo diferencial”, esta serie permite aproximaciones a funciones, la idea principal es poder aproximar los valores de una función $f(x)$ para cualquier punto x , partiendo de un polinomio basado en una serie de potencias infinita.

La expansión de la serie de Taylor puede ser dada en términos de una sumatoria infinita, y los coeficientes son las sucesivas derivadas de una función f diferenciable infinitas veces en $x = a$.

$$f(x) = \sum_{n=0}^{\infty} \frac{f^n(a)}{n!} (x - a)^n$$

Se puede fácilmente deducir esta fórmula mediante el uso de la coinducción. Sin embargo, para esto, se requiere de una operación c la cual permita

calcular la cabeza y la cola del polinomio infinito.

Entonces sea $c(x)$ una coálgebra que actúa de la siguiente forma:

I. $f^n(a)$

II. $f^n(x)$, n es el número de derivadas en ambas operaciones .

La primera me permite observar la cabeza y la segunda me permite observar la cola o el polinomio infinito.

Con la siguiente expansión de series de potencias o polinomio de series infinitas, se busca llegar a la serie de Taylor.

$$f(x) = c_0 + c_1(x - a) + c_2(x - a)^2 + c_3(x - a)^3 + \dots + c_n(x - a)^n$$

Entonces al realizar la primera operación de la coálgebra, evaluando a en la función anterior, se obtiene:

$$f(a) = c_0$$

Si se calcula la primer derivada o se aplica la operación de la función se obtiene lo siguiente:.

$$f'(x) : c_1 + 2c_2(x - a) + 3c_3(x - a)^2 + \dots$$

Con este proceso se ha calculado la primer cabeza que es c_0 y su cola infinita que es $f'(x)$. Se continua con las operaciones de c , evaluando a en $f'(x)$ y calculando la segunda derivada , se obtiene los siguiente:

$$f'(a) = c_1$$

$$f''(x) : 2c_2 + 3 \cdot 2c_3(x - a) + \dots$$

Evalutando a en la segunda derivada de $f(x)$

$$f''(a) = 2c_2 \text{ entonces } c_2 = \frac{f''(a)}{2}$$

Este proceso se puede hacer indefinidamente y al continuar con el proceso, de evaluar a en f^n y calculando la n derivada de $f(x)$, se puede observar que:

$$f^{(n)}(a) = n!c_n$$

De esta forma se obtiene la serie de Taylor.

$$f(x) = \sum_{n=0}^{\infty} \frac{f^n(a)}{n!} (x - a)^n$$

$$f(a) + f'(a)(x - a) + \frac{f''(a)}{2}(x - a)^2 + \frac{f'''(a)}{3!}(x - a)^3 + \dots$$

La serie de Taylor proporciona una buena forma de aproximar el valor de una función en un punto y sus derivadas en otro punto.

Uniendo las derivadas de f se obtiene el conjunto:

$$\mathbb{A} = f(x), f'(x), f''(x), f'''(x) \dots f^n(x)$$

Si $R^{<\omega}$ es el conjunto de secuencias de coeficientes de Taylor, es decir, de $\alpha \in R^{<\omega}$ o $\sum \frac{\alpha_i}{i!} x^i < \infty$ para $0 < x$

$$T\{f\} = [f(0), f'(0), f''(0)..]$$

$$\tilde{T}\{\alpha\} = \sum_{i=0}^{\infty} \frac{\alpha_i}{i!} x^i$$

$$T = \tilde{T}$$

Con las operaciones anteriores se observa la identidad del conjunto $\mathbb{A}$ con $\sum \frac{\alpha_i}{i!} x^i < \infty$. Obteniendo el morfismo identidad T

$$A \xrightarrow{T} R^{<\omega}$$

Con las dos operaciones plateadas de c $f^n(a)$ y $f^n(x)$, donde el primero calcula la cabeza O y el segundo calcula la cola D , entonces se tiene un morfismo, el cual parte desde un elemento del conjunto $\mathbb{A}$ y se llega a una pareja ordenada (O, D) perteneciente al producto $(\mathbb{R} \times \mathbb{A})$.

$$\mathbb{A} \xrightarrow{(O,D)} \mathbb{R} \times \mathbb{A}$$

Por definición de coálgebra final (los números reales) existe un único morfismo tal que:

$$\mathbb{R}^\omega \xrightarrow{(head,tail)} \mathbb{R} \times \mathbb{R}^\omega$$

.

Si el funtor identidad $i : \mathbb{R} \longrightarrow \mathbb{R}$ y si $A \xrightarrow{T} R^{<\omega}$ al hacer el producto entre ambos funtores se obtiene:

$$\mathbb{A} \times \mathbb{R} \xrightarrow{(\mathbb{R} \times T)} \mathbb{R} \times \mathbb{R}^\omega$$

Con lo anterior se puede visualizar la serie de Taylor por medio de la siguiente estructura:

$$\begin{array}{ccc} \mathbb{A} & \xrightarrow{(O,D)} & \mathbb{R} \times \mathbb{A} \\ \downarrow T & & \downarrow (\mathbb{R} \times T) \\ R^{<\omega} & \xrightarrow{(head,tail)} & \mathbb{R} \times \mathbb{R}^\omega \end{array}$$

De esta forma la serie de Taylor puede ser abordada por medio de la coinducción

Ejemplo con la serie de Maclaurin.

Tomando $f(x) = \text{sen}(x)$, formar la serie de Maclaurin. Esta serie es un caso particular de la serie de Taylor pero centrado o evaluado en cero, entonces se parte desde la siguiente sumatoria:

$$f(x) = \sum_{n=0}^{\infty} \frac{f^n(0)}{n!} x^n = f(0) + f'(0)x + \frac{f''(0)}{2!}x^2 + \frac{f^3(0)}{3!}x^3 + \frac{f^4(0)}{4!}x^4 \dots$$

Entonces sea $c(x)$ una coálgebra que actúa de la siguiente forma:

I. $f^n(0)$

II. $f^n(x)$, n es el número de derivadas en ambas operaciones

Al aplicar 1 y 2 se obtiene la cabeza y la cola respectivamente. Derivando varias veces $f(x)$ y luego evaluando en 0 en cada derivada n -ésima. Se tiene la siguiente lista

$$n \quad f^n(x) = \text{cola} \quad f^n(0) = \text{cabeza}$$

$$0 \quad f(x) = \text{sen}(x) \quad f(0) = 0$$

$$1 \quad f'(x) = \text{cos}(x) \quad f'(0) = 1$$

$$2 \quad f''(x) = -\text{sen}(x) \quad f''(0) = 0$$

$$3 \quad f^3(x) = -\text{cos}(x) \quad f^3(0) = -1$$

$$4 \quad f^4(x) = \text{sen}(x) \quad f^4(0) = 0$$

$$5 \quad f^5(x) = \text{cos}(x) \quad f^5(0) = -1$$

Así sucesivamente, la secuencia se repite tras la tercera derivación, así que la serie de potencias pedida es:

$$\sum_{n=0}^{\infty} \frac{f^n(0)}{n!} (x)^n = f(0) + f'(0)x + \frac{f''(0)}{2!}x^2 + \frac{f^3(0)}{3!}x^3 + \frac{f^4(0)}{4!}x^4 \dots$$

$$\sum_{n=0}^{\infty} \frac{(-1)^n x^{2n+1}}{(2n+1)!} = x - \frac{x^3}{3!} - \frac{x^5}{5!} - \frac{x^7}{7!}$$

Con lo anterior se puede concluir que el teorema fundamental del cálculo, la serie de Taylor y de Maclaurin, estudian a los números reales y por ende se trabaja con un proceso de coinducción, dado que son parte de una coálgebra final, como son los números reales. Sin embargo existen muchos conceptos que estudian a $\mathbb{R}$, como los límites, las sucesiones, los ejemplos particulares

de funciones y series, entre otros; los cuales no se abordaron en este trabajo puesto que es un trabajo extensivo y se deja postergado como un trabajo futuro.

En este capítulo se abordó la dualidad que existe entre los números naturales y los números reales, por medio de la teoría de categorías, esta teoría permite estudiar la estructura de los cuerpos matemáticos y se dejan a un lado los conceptos de pertenencia de los elementos de un conjunto. Los morfismos son funciones que van desde un objeto matemático hasta otro objeto de la misma categoría. Con ayuda de los funtores se puede relacionar dos categorías, esto se hizo con el fin de caracterizar álgebras y cóálgebras finales.

Presentando la dualidad, se observaron los siguientes resultados:

- Los números naturales son álgebras de tipo $0 - 1$ porque tienen como operación unaria al sucesor y como operación cero-aria al elemento distinguido cero. Además, los números naturales se establecen como un objeto inicial dado que este objeto se construye inductivamente. Al ser un objeto inicial, los números naturales cumplen con dos principios; el de inducción que es el método de demostración y el de recursión el cual establece que las propiedades y operaciones estén bien definidas.

- Los números reales son cóálgebras finales, que forman parejas (U, c) en el cual U es un conjunto y c es una función, tal que: $c : U \longrightarrow T(U)$. En las cóálgebras no se construyen como en las álgebras, si no que se observan o analizan los elementos del conjunto de llegada de $T(U)$, con el fin de conocer las operaciones que actúan en el dominio U , además, las cóálgebras también se pueden operar o componer, esto implica que si una cóálgebra es final es porque existe un único morfismo que va desde una cóálgebra arbitraria a la cóálgebra final. Según [17] y [13] los números reales son la cóálgebra final, dado que se puede estudiar la representación decimal del intervalo $[0, 1)$ con ayuda de los operadores de observación *head* y *tail*, formando el morfismo único (h, t) que va de cualquier cóálgebra arbitraria a la cóálgebra final. Como los números reales son un objeto final entonces cumplen con el principio de coinducción y el de corecursión.

Conclusiones

- A lo largo de este trabajo se han estudiado cuatro teorías, con el fin de comprender la coinducción matemática y ver desde otra perspectiva a los números reales. La primera de ellas es la teoría de conjuntos que permite ver de forma clara y precisa los principios de inducción y recursión; sin embargo, no es suficiente para entender el concepto de coinducción. Por esta razón, se estudió la teoría de puntos fijos que es una rama de la topología la cual va más allá de la teoría de conjuntos. Esta teoría permite evidenciar la dualidad que existe entre la inducción y la coinducción mediante los puntos prefijos y puntos postfijos respectivamente, además permite definir los conjuntos inductivos y coinductivos. Con la ayuda de la teoría de retículos se pudo identificar a los conjuntos inductivos como la intersección de todos los puntos prefijos; es decir, la existencia de un ínfimo o menor punto prefijo, y dualmente se puede definir a los conjuntos coinductivos como la unión de todos los puntos postfijos. Es decir, la existencia de un supremo o mayor punto postfijo. Entonces con estos resultados se puede comprender que la coinducción es dual a la inducción, lo que implica que los números reales son duales a los naturales. Esta dualidad se abordó mediante la teoría de categorías la cual permite estudiar la estructura de los objetos, es por ello que se definen a los números naturales como álgebras iniciales mientras que a los números reales como coálgebras finales, en consecuencia se define a $\mathbb{N}$ como objeto inicial y a $\mathbb{R}$ como objeto final.
- Los números naturales son álgebras iniciales, puesto que existe un único morfismo que relaciona el álgebra inicial con un álgebra arbitraria. Dualmente, los números reales son una coálgebra final porque existe un único morfismo que relaciona una coálgebra arbitraria con la final.

- Los números naturales son objeto inicial puesto que cumplen con la característica de ser la intersección de todos los conjuntos inductivos. Esto permite establecer los principios de inducción y recursión, el primero es un método de demostración de propiedades mientras que el segundo es un método que garantiza la buena definición de las operaciones entre los elementos de $\mathbb{N}$. La inducción es el paso básico para el desarrollo de la aritmética elemental.
- Los números reales son objetos finales puesto que cumplen con las características de una coálgebra final. Al ser un objeto final se establecen los principios de coinducción y de corecursión. El primero es un método de demostración que emplea la observación y el segundo permite establecer que las operaciones entre listas infinitas estén bien definidas, y el paso básico para el desarrollo del cálculo elemental es el principio coinductivo que se halla oculto en cada análisis matemático.
- Los números reales pueden ser tratados con el método coinductivo porque pertenece a las matemáticas basadas en el análisis y no a las matemáticas inductivas, como es el caso de los números naturales, los cuales se pueden construir recursivamente mediante la operación sucesor. Es por ello que consideramos de manera ingenua que existen obstáculos de aprendizaje en los estudiantes a la hora de abordar conceptos relacionados con los números reales como los vistos en el cálculo, puesto que estos conceptos necesitan del análisis para ser desarrollados. Como trabajo futuro se propone estudiar cómo el principio de coinducción se podría empezar a evidenciar en la educación básica, media y universitaria.
- Este trabajo se inició con los fundamentos que proporcionan los cursos del programa de licenciatura de educación básica con énfasis en matemáticas como lo son Teoría de Conjuntos, Álgebra Moderna, Estructuralismo de las Matemáticas del siglo XIX y XX y algunos otros que hacen parte de la línea de historia. Sin embargo, a lo largo de esta investigación histórica se encontró que se puede partir no solo de ZFC sino de las teorías axiomáticas consistentes que no incluyen el axioma de fundamentación, como lo es la teoría de hiperconjuntos de Peter Aczel, que consiste en permitir la existencia de conjuntos que formen cadenas de pertenencia con descenso infinito. Esta teoría dio inicio al concepto de coálgebra y coinducción que son duales a las álgebras y la

inducción. Esto implica que se puede partir de la teoría de Peter Aczel y hacer un estudio tal vez más detallado de este trabajo.

- Con los ejemplos de la series de Taylor y el Teorema Fundamental del Cálculo, se ha ilustrado la utilización coinductiva en el análisis de $\mathbb{R}$. En consecuencia conocer y comprender el concepto de coinducción es importante para la formación de los estudiantes, puesto que es fundamental para el estudio de los conceptos que involucran los números reales y para desarrollar un pensamiento analítico.
- Se evidencia que existe una dualidad entre $\mathbb{N}$ y $\mathbb{R}$, porque el primero es un conjunto inductivo, mientras que el segundo es un conjunto coinductivo. La inducción permite construir mientras que la coinducción permite deconstruir; la inducción se encuentra inmersa en la aritmética y de forma dual la coinducción se encuentra inmersa en el análisis. Además, se puede asociar la inducción con el método sintético, el cual es un proceso de razonamiento que tiende a reconstruir un todo, a partir de los elementos distinguidos. Por otro lado, la coinducción se puede asociar con el método analítico, pues que consiste en la deconstrucción de un todo, descomponiéndolo en sus partes para observar las causas, la naturaleza y los efectos.

Bibliografía

- [1] Aczel, P. (1988). Non-well-founded sets. Stanford: Library of Congress Cataloging in Publication Data.
- [2] Balbes, R., & Dwinger, P.(1974)*Distributive Lattices*,Columbia: University of Missouri Press.
- [3] Blyth, T. (2005) *Lattices and Ordered Algebraic Structures*. San Francisco: Springer.
- [4] Cejas, A. M. (2000). Teoría de Punto Fijo. Las matemáticas del siglo XX una mira en 1001 artículos. p,(329-332).
- [5] Crespín,D.(2014). *Números Naturales y Recursión*. Recuperado el 13 de Agosto de 2014, <http://www.matematica.ciens.ucv.ve/dcrespin/Pub/NumNat.pdf>
- [6] Dedekind, R. (1998). *¿Qué son y para qué sirven los números?* Madrid: Alianza. Primera versión en alemán, 1888.
- [7] González, L. D. (2007). *Coinducción: de la Teoría de Categorías a la Programación Funcional (tesis de pregrado)*. Universidad Nacional Autónoma, Mexico.
- [8] García, J. (2005-2006). *Conjuntos Ordenados. Retículos y álgebras de Boole*. Recuperado el 21 de Marzo del 2014, de [http://www.ugr.es/jesusgm/Curso2005-2006/Matematicas Discreta/Ordenes.pdf](http://www.ugr.es/jesusgm/Curso2005-2006/Matematicas%20Discreta/Ordenes.pdf).
- [9] Geiss, C.(2005) *Algunas Propiedades De Los Números Los Naturales*. Recuperado el 7 Marzo de 2014, en [http://www. matem.unam.mx/christof/ cursos/AS/Nota2.pdf](http://www.matem.unam.mx/christof/cursos/AS/Nota2.pdf)

- [10] Gómez, J.L. Un paseo alrededor de la teoría de conjuntos. *LA GACETA DE LA RSME*. Vol.(11), pp.(45-96).
- [11] Hilton, P.(1982). *Curso de álgebra moderna*(pp. 71-80). España: EDITORIAL REVERTÉ, S.A.
- [12] Hrbacek, K., & Jech, T. (1999). *Introduction to set theory*, New York: Marcel Dekker.
- [13] Jacobs, B. & Rutten, J.(1997). A Tutorial on (Co)Algebras and (Co)Induction. *EATCS Bulletin*. Vol.(62), pp.(62-212).
- [14] Jacobs, B. & Rutten, J. (1997) A Tutorial on (Co)Algebras and (Co)induction*. *EATCS Bulletin*. (62), pp. (222-259).
- [15] Munkres, J.R.(2000). Segunda Edición, *Topología*, Madrid: Isabel Campella, pp.(59-85)
- [16] Ortiz, G. R.(2010). *Borrador de una Propuesta para un curso de Teoría de Conjuntos y Sistema Numéricos*. Santiago de Cali: Universidad del Valle.
- [17] Ortiz, G., & Valencia, S. (2010). La categoricidad de los reales en Hilbert. *Revista Brasileira de História da Matemática*. Vol. 10, No. 19, pp.(39-65).
- [18] Pavlović, D. and Escardo, M. (1998) Calculus in coinductive form. In: *Proceedings of the 13th Annual IEEE Symposium on Logic in Computer Science*, IEEE Computer Society Press, pp. (408–417).
- [19] Pavlović, D., & Pratt, V. (2002, 09). The continuum as a final coalgebra. *Theoretical Computer Science*, volumen(280), (105-106).
- [20] Rafael, F., Isaacs, G., Sabogal, P., & Sonia, M.(2005). *Números naturales y recursividad*. Recuperado el 30 de Enero del 2014, de <http://matematicas.uis.edu.co/risaacs/AMA/doc/Recursivas.pdf>
- [21] Roman, S.(2008). *Lattices and ordenred set*, Newyork USA: Springer, pp.(263-269)
- [22] Serrano, F.(2008). *Elementos de matemáticas formalizadas en Isabelle/Isar* (Tesis doctoral). Universidad de Sevilla, Sevilla, España.

- [23] Téllez, A.(2011).*Existencia de coálgebras finales para funtores polinomiales*(Tesis de pregrado).Universidad Del Valle, Santiago de Cali.
- [24] Zalamea, F. (2009). Hacia una filosofía sintética de las matemáticas contemporáneas. Bogotá: Universidad Nacional,pp.(21-33)
- [25] Zermelo, E. (1908). Investigations in the foundations of set theory I. En H. Ebbinghaus, C. Fraser, & A. Kanamori (Eds), Ernest Zermelo. Collected Works (pp. 189-228). Berlin Freiburg: Springer-Verlag. (También en Van Heijenoort (1967, pp. 199-215)).
- [26] Zermelo, E. (1930). On boundary numbers and domains of sets. New investigations in the foundations of set theory. En H. Ebbinghaus, C. Fraser, & A. Kanamori (Eds), Ernest Zermelo. Collected Works (pp. 401-431). Berlin Freiburg: Springer-Verlag.