



Aquellas pequeñas cosas

*Entre infinitesimales e intuiciones:
una introducción al análisis no estándar*

Matías Nicolás Sollier

Facultad de Ingeniería Química
Universidad Nacional del Litoral

Agosto 2025



Concurso de monografías UMA - 2025



*“Uno se cree que las mató el tiempo y la ausencia
Pero su tren vendió boleto de ida y vuelta
Son aquellas pequeñas cosas
Que nos dejó un tiempo de rosas”*

— Joan Manuel Serrat, *Aquellas pequeñas cosas*

Prefacio

“Yo creo en muchas cosas que no he visto, y ustedes también, lo sé.”

Con estas palabras, Willie Colón da comienzo a su versión de la canción *Oh, qué será?*, en una introducción inspirada en las palabras de la escritora Clarice Lispector. De la misma manera, me gustaría comenzar esta monografía.

Los objetos de estudio de las ciencias naturales —como la física, la química o la astronomía— resultan ser fenómenos o realidades tangibles, que se pueden experimentar por medio de los sentidos, que se pueden “ver”. El conocimiento en ellas se genera, en primer medida, a partir de experiencias empíricas. A pesar de esto, a la hora del desarrollo de teorías científicas, es muy común definir entidades hipotéticas que sirvan para modelar la realidad, asumiéndose su existencia *a priori*. Por lo tanto, no resulta descabellado pensar en que un científico de esta área “crea en cosas que no ha visto”. La matemática, por otro lado, no trabaja con objetos que puedan “verse”, en el sentido usual. Sin embargo, considero que la frase que da inicio a esta monografía describe de manera muy acertada el pensamiento de la comunidad matemática sobre el eje central de la misma: los *infinitesimales*.

Estoy seguro que toda persona que haya tomado un curso de cálculo trató de imaginarse cantidades infinitesimales. Intentar entender como luciría en la recta numérica un número mayor que cero pero menor que cualquier otro número real positivo es un buen ejercicio mental. A medida que uno avanza en su formación académica como estudiante de matemática —pasando por el cálculo en varias variables, análisis en espacios métricos, análisis funcional, entre otros— incorpora el concepto de límite y las demostraciones utilizando argumentos ε - δ , dejando de lado las dubitaciones que podría llegar a tener sobre aquellas pequeñas cosas, pues el formalismo actual funciona muy bien y resulta lo suficientemente intuitivo.

Cuando me enteré de la existencia de una definición rigurosa —a los ojos de la matemática actual— para los infinitesimales, me invadió la curiosidad. Grande fue mi sorpresa al ver la cantidad de fundamentación lógica y algebraica que hay detrás de la mismo. Sin importar la intuición que tuvieran Newton y Leibniz, les habría sido imposible obtener algo similar utilizando las herramientas matemáticas de su momento.

Escribo esta monografía con el fin de poder brindar un breve resumen de mi trabajo de investigación en este tema. Traté de presentar la mínima cantidad de contenido teórico necesario para poder obtener resultados interesantes y bonitos. Como se menciona más adelante en la introducción, el objetivo principal es que el lector pueda interpretar intuitivamente las conclusiones, a medida que vaya leyendo, así como reflexionar sobre las mismas.

Para finalizar estas palabras, me gustaría agradecer a mi novia Camila Roncoroni que siempre me motivó y me ayudó a salir adelante en los momentos donde sentía que no podía avanzar con la escritura de este trabajo; a mi amigo Tomás Gerbazoni con quién comenzamos a investigar juntos sobre estos temas; a mis profesores Marcelo Actis y Pablo Quijano que me ayudaron y guiaron en toda esta nueva experiencia; y finalmente a mi familia, a mi papá, a mi mamá y a mi hermano, por haberme apoyado siempre.

Índice General

Prefacio	II
1 Introducción	1
1.1 Sobre el concepto de infinitesimal	1
1.2 Sobre la monografía	2
2 Construcción	3
2.1 Generando una intuición	3
2.2 Caracterizando la grandeza	5
2.3 Manos a la obra	6
2.4 ¿Nuevos elementos?	8
2.5 Clasificación y propiedades	10
2.6 El principio de transferencia	13
3 Análisis no estándar	14
3.1 Sucesiones	14
3.2 Continuidad	17
3.3 Diferenciación	19
3.4 Integración	24
4 Aplicaciones	29
4.1 Ecuaciones diferenciales ordinarias	29
4.2 La ecuación de onda en una cuerda	31
4.3 Otras aplicaciones	33
5 Epílogo	35
Bibliografía	36

1. Introducción

“Y... y no olvidar que la estructura del átomo no es percibida aunque se sepa que existe. Sé de muchas cosas que no vi. Y ustedes también. No se puede dar prueba de la existencia de lo que es más verdadero, la cosa es creer. Creer llorando.”

— Clarice Lispector, *La hora de la estrella*

1.1. Sobre el concepto de infinitesimal

Seguramente el término *infinitesimal* o *infinitésimo* no resulta para nada ajeno al lector. Es comúnmente utilizado para describir una cantidad positiva que es muy pequeña, más pequeña de lo que se pueda imaginar, más incluso que cualquier número real positivo, pero no es nula (es decir, no es cero).

Originalmente fueron Newton y Leibniz quienes acuñaron este concepto, apoyándose en él para justificar una gran cantidad de las deducciones y los razonamientos en sus escritos. Si bien es indiscutible que la semilla que plantaron Newton y Leibniz dio uno de los frutos más importantes de la matemática, como lo es el *cálculo infinitesimal*, la falta de rigurosidad y formalidad no resultaría del agrado de los matemáticos venideros.

Posteriormente el cálculo tendría una reversión más adecuada a los ojos del rigor. Los trabajos de Cauchy, Riemann y Weierstrass permitieron formalizar los resultados de esta área utilizando un concepto denominado *límite*. Esta manera de entender el cálculo se volvió estándar a lo largo del mundo, y sigue vigente hasta el día de hoy. Más aún, el concepto de límite es clave en el *análisis matemático*, otra rama de la matemática que continuaría el legado del *cálculo* estudiando sus objetos con un mayor nivel de abstracción y generalidad.

A pesar de la actual hegemonía del límite, los infinitesimales se encuentran muy lejos de donde habita el olvido. Es muy común encontrarlos en libros de física a la hora de deducir ecuaciones diferenciales o al hacer referencia a “un pequeño desplazamiento”, “una pequeña cantidad de tiempo”, o al dividir un cuerpo sólido en “regiones de volumen muy pequeño”, entre muchos otros usos.

Pero entonces, si el valor intuitivo de los infinitesimales es tan apreciado, ¿existe una forma de describir matemáticamente cantidades infinitesimales? La respuesta a esta pregunta es sí, sí existe. A mediados del siglo XX el matemático Abraham Robinson publicó un libro titulado “*Non-standard Analysis*” [8]. En este demostró la existencia de objetos matemáticos que representen cantidades infinitesimales, y cómo obtener muchos de los resultados modernos de análisis en términos de infinitesimales.

El análisis no estándar, NSA (por sus siglas en inglés), brindó una nueva forma de ver el análisis, y por consecuencia el cálculo. Una forma que intenta contemplar las intuiciones de los infinitesimales, sin perder el formalismo lógico que caracteriza a la matemática. Hasta el día de la fecha se han publicado varios libros y trabajos sobre este tema. En 1976 H. J. Keisler publicó “*Elementary Calculus: An Infinitesimal Approach*” [6], un curso básico de cálculo, pero fundado sobre el análisis no estándar.

Aquel lector con un interés más profundo en la historia del cálculo puede consultar el libro de Edwards [3].

1.2. Sobre la monografía

El objetivo de esta monografía es ofrecer una visión general del análisis no estándar. No se busca entrar en tanto detalle ni en el rigor de la construcción de los objetos, sino más bien generar la intuición de por qué todas las deducciones, demostraciones y teoremas funcionan y tienen sentido.

El Capítulo 2 aborda la construcción del conjunto de los números hiperreales, siguiendo de cerca el enfoque de Robert Goldblatt [4]. Allí se introducen los conceptos y definiciones fundamentales que servirán como base para el resto de la monografía, poniendo énfasis en la comprensión de la estructura y las propiedades esenciales de los hiperreales. El objetivo de este capítulo no es convencer al lector de que una forma es más conveniente que la otra, sino ofrecer una perspectiva distinta, mostrando cómo se pueden derivar teoremas del análisis estándar usando el lenguaje y las herramientas del análisis no estándar.

El Capítulo 3 desarrolla aspectos propios del análisis no estándar. Se seguirá utilizando el libro de Goldblatt [4] en combinación con el libro de Keisler [6] para brindar una interpretación más intuitiva.

Finalmente, el Capítulo 4 se centra en mostrar aplicaciones de las ideas trabajadas en capítulos anteriores, particularmente en el contexto de ecuaciones diferenciales y en problemas de física, con el fin de ilustrar el alcance y la utilidad de esta teoría.

2. Construcción

“Subiu a construção como se fosse máquina.”

— Chico Buarque, *construção*

La contribución de Abraham Robinson como el creador del análisis no estándar se basó en demostrar rigurosamente la existencia de una extensión lógica de los números reales que incluye cantidades infinitesimales (entre otras), y que además conserva todas las propiedades lógicas del análisis “estándar”, por medio del principio de transferencia. Si bien este resultado fue muy importante, Robinson no ofreció una construcción explícita de este nuevo conjunto. Fueron W.A.J. Luxemburg y H.J. Keisler quienes desarrollaron e implementaron una construcción de este, al que denominaron el conjunto de los números hiperreales, por medio de un ultraproducto. En este capítulo adoptaremos precisamente este enfoque para construir los números hiperreales.

2.1. Generando una intuición

Comencemos pensando en cómo podríamos representar una cantidad infinitesimal. Claramente no nos basta con utilizar un número real, ya que partir de que existe un número real mayor que cero y menor que cualquier número real positivo resulta en una contradicción, por lo que tenemos que pensar por fuera de este conjunto.

Usemos ahora un enfoque más dinámico. Pensemos en una variable que puede tomar distintos valores reales positivos. Llamemos a esta variable x y a su n -ésimo valor x_n . Si se tiene que los valores x_n decrecen indefinidamente de forma de ser más pequeño que cualquier otro número real positivo dado, entonces podríamos interpretar esta cantidad como una cantidad infinitesimal, pues nunca toma el valor cero y “eventualmente” se hará tan pequeña como se quiera. De hecho, esta fue la definición que dio Cauchy en el capítulo 2 de su *Cours d’analyse* [2]. Él llamo a estas cantidades *infinitamente pequeñas*. Un ejemplo de esto podría ser una variable que tome por valores $1, \frac{1}{2}, \frac{1}{3}, \dots$. A la luz de los conceptos actuales de análisis podríamos afirmar que la variable x es infinitamente pequeña si la sucesión de valores que $\langle x_1, x_2, x_3, \dots \rangle$ es decreciente y converge a 0.

Estas variables infinitamente pequeñas son muy similares a las cantidades infinitesimales que deseamos representar. Avancemos un poco más con esta definición.

Pensemos ahora en dos variables x e y que toman los valores

$$1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots \quad \text{y} \quad \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \dots$$

respectivamente¹. Ambas variables son infinitamente pequeñas acorde a la definición. Pero analizándolas no sería descabellado pensar que la segunda variable “se hace pequeña más rápido” que la primera variable. Una forma de justificar este razonamiento es ver que, para cualquier índice que elijamos, el valor que toma la segunda variable es menor que el valor que toma la primera. Simbólicamente lo que estamos queriendo decir es que $x_n > y_n$ para todo $n \in \mathbb{N}$. Esto nos abre la puerta para pensar en cantidades infinitesimales más pequeñas que otras.

¹Los términos de las sucesiones son $x_n = n^{-1}$ y $y_n = 2^{-n}$ respectivamente.

Introducimos ahora una nueva variable z que toma valores

$$2, \frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \dots$$

Si bien es falsa la afirmación que $x_n > z_n$ para todo $n \in \mathbb{N}$, pues $x_1 < z_1$, sería muy contraintuitivo pensar que z no es más pequeña que x , pero y si lo es. Todo esto se justifica en que z e y solo difieren en un valor, y diferir un solo valor no es suficiente comparado con coincidir en todos los números naturales. Lo que nos lleva a pensar en que, si bien z e y son variables diferentes, al coincidir en muchos de los valores que toman, la cantidad infinitesimal que representan podría ser la misma.

Todas estas afirmaciones y deducciones intuitivas que estamos realizando en el aire nos llevan a la conclusión de que debemos construir una forma de comparar sucesiones. De esta manera podremos afirmar con certeza cuándo dos variables representan a la misma cantidad infinitesimal, o si una representa una cantidad infinitesimal más grande que la otra. Además, como las variables vienen biunívocamente determinadas por la sucesión de valores que toman, a partir de ahora trabajaremos únicamente con las sucesiones, dejando a las variables de lado.

Para realizar esto introducimos una notación que utilizaremos muy frecuentemente a lo largo de esta primera parte de la monografía. Si tenemos las sucesiones

$$r = \langle r_1, r_2, \dots \rangle \quad \text{y} \quad s = \langle s_1, s_2, \dots \rangle,$$

entonces definimos el *conjunto de coincidencias* de r y s como

$$\llbracket r = s \rrbracket := \{n \in \mathbb{N} : r_n = s_n\}.$$

Es decir, tal como indica su nombre, el conjunto de los índices en los cuales los términos de las sucesiones coinciden. Por ejemplo, si tenemos

$$r = \left\langle 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots \right\rangle, \quad s = \left\langle 1, \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \dots \right\rangle \quad t = \left\langle 0, \frac{1}{2}, 0, \frac{1}{4}, \dots \right\rangle,$$

entonces

$$\llbracket r = s \rrbracket = \{1, 2\}, \quad \llbracket s = t \rrbracket = \{2\}, \quad \llbracket r = t \rrbracket = \{2, 4, 6, \dots\}.$$

Es importante no confundir este conjunto con “los elementos que tienen en común” las sucesiones. Por ejemplo, las sucesiones $x = \langle 0, 1, 0, 1, \dots \rangle$, $y = \langle 1, 0, 1, 0, \dots \rangle$ tienen los mismos elementos como conjuntos, pero $\llbracket x = y \rrbracket = \emptyset$.

De manera análoga se puede definir el conjunto $\llbracket r < s \rrbracket := \{n \in \mathbb{N} : r_n < s_n\}$. En general, para cualquier proposición P que tenga como variable dos números reales, se puede pensar en el conjunto $\llbracket P(r, s) \rrbracket := \{n \in \mathbb{N} : P(r_n, s_n) \text{ es verdadero}\}$. Por ejemplo $\llbracket r \neq s \rrbracket$ o $\llbracket r \geq s \rrbracket$.

Volvamos a nuestro conjunto de coincidencias. Podemos interpretar a $\llbracket r = s \rrbracket$ como un indicador o una medida de qué tan cierto es “ $r = s$ ”. De forma que, si $\llbracket r = s \rrbracket = \mathbb{N}$ se tiene que las sucesiones coinciden en todos sus términos, y si $\llbracket r = s \rrbracket = \emptyset$ no coinciden en ninguno. En estos casos extremos es donde sabemos exactamente si las sucesiones representan o no la misma cantidad infinitesimal (sí en el primer caso, y no en el segundo). El problema radica en cómo interpretamos otros valores del conjunto de coincidencias. Por

ejemplo, qué pasa si el conjunto de coincidencias es finito, o si es de la forma $\{2, 4, 6, \dots\}$. Si $\llbracket r = s \rrbracket = \{2\}$, como en el ejemplo anterior, no tendría mucho sentido pensar en que r y s representan lo mismo, ya que solo coinciden en un término. Lo mismo se podría pensar de $\llbracket r = s \rrbracket = \{1, 2\}$. En general, si el conjunto de coincidencias es finito, no nos gustaría afirmar que las sucesiones representan la misma cantidad, pues difieren en muchos más términos de los que coinciden. Por otro lado, si $\llbracket r = s \rrbracket = \{2, 3, \dots\} = \mathbb{N} - \{1\}$, siguiendo la misma lógica que antes, no sería de nuestro agrado afirmar que r y s no representan la misma cantidad, pues coinciden en muchos más términos de los que difieren. Los conjuntos tales que su complemento respecto a $\mathbb{N}$ es finito, es decir conjuntos de la forma $\mathbb{N} - A$ con A finito, se denominan *cofinitos*. Gracias a nuestro análisis previo podemos concluir que, si nuestro conjunto de coincidencias es finito, entonces queremos que las sucesiones no representen la misma cantidad. En cambio, si nuestro conjunto de coincidencias es cofinito, las sucesiones representarán la misma cantidad.

Debajo de todo esto subyace la idea de que

“dos sucesiones representarán la misma cantidad
si coinciden en una cantidad grande de términos”,

lo que se traduce en decir que

“su conjunto de coincidencias es grande”.

Pero ¿qué significa grande? Es difícil de poner en palabras, pero ya sabemos ejemplos de conjuntos que no queremos que califiquen como grandes (el conjunto vacío y los conjuntos finitos), y de conjuntos que sí queremos ($\mathbb{N}$ y los conjuntos cofinitos). ¿Qué hay del resto de posibles conjuntos de coincidencias? Ya que podrían ser cualquier subconjunto de $\mathbb{N}$, ¿cómo hacemos para *filtrar* los que son grandes de los que no?

2.2. Caracterizando la grandeza

*“¿De qué tamaño es tu amor?
Que es muy grande para mí.
Me pregunto cada vez.”*

— Héctor Lavoe,
¿De Qué Tamaño Es Tu Amor?

Aquí es donde entra en juego uno de los conceptos más importantes de la construcción de los números hiperreales: el *filtro*.

Dado un conjunto I no vacío, definimos un filtro sobre I como una colección no vacía $\mathcal{F} \subseteq \mathcal{P}(I)$ que cumple con las siguientes propiedades:

- (1) Si $A, B \in \mathcal{F}$, entonces $A \cap B \in \mathcal{F}$.
- (2) Si $A \in \mathcal{F}$ y $A \subseteq B \subseteq I$, entonces $B \in \mathcal{F}$.

Un ejemplo de filtro sobre un conjunto finito $\{1, 2, 3, 4\}$ es

$$\mathcal{F} = \{\{3, 4\}, \{1, 3, 4\}, \{2, 3, 4\}, \{1, 2, 3, 4\}\}.$$

Dado cualquier conjunto no vacío I una forma rápida de crearse un filtro sobre I es definir para un subconjunto $A \subseteq I$ la colección $\{B \subseteq I : A \subseteq B\}$. Es sencillo probar que esta colección resulta ser un filtro sobre I . De hecho, este filtro se denomina el *filtro generado por A* . Además, decimos que un filtro es *propio* si no contiene al conjunto vacío, pues si $\emptyset \in \mathcal{F}$, entonces $\mathcal{F} = \mathcal{P}(I)$.

Utilizaremos los filtros sobre $\mathbb{N}$ para caracterizar los conjuntos que consideraremos “grandes”. Es decir, $A \in \mathcal{F}$ si y solo si “ A es grande”. Notemos que ambas propiedades de los filtros son deseables a la hora de representar conjuntos grandes. La primera indica que si tenemos dos conjuntos grandes, entonces comparten una cantidad grande de elementos. La segunda es todavía más intuitiva, si un conjunto contiene a un conjunto grande, entonces este debe ser grande también.

A simple vista parece que cualquier filtro nos basta, pero hay un problema. Si miramos el ejemplo anterior del filtro $\mathcal{F}$ sobre el conjunto $\{1, 2, 3, 4\}$, sucede que $\{4\} \notin \mathcal{F}$ y $\{1, 2, 3\} \notin \mathcal{F}$. Es decir, existe un subconjunto de $\{1, 2, 3, 4\}$ tal que ni este ni su complemento pertenecen al filtro. Si bien este es un ejemplo para un conjunto finito, esto puede ocurrir para un filtro sobre $\mathbb{N}$ (considerar el filtro generado por $\{1, 2\}$ y el conjunto $\{1\}$). Recordemos que bajo nuestra lectura del conjunto de coincidencias, si r y s representan la misma cantidad, entonces $\llbracket r = s \rrbracket$ es grande, y si no representan la misma cantidad, entonces difieren en una cantidad grande de términos, es decir, $\llbracket r \neq s \rrbracket = \mathbb{N} - \llbracket r = s \rrbracket$ es grande.

Para solucionar esto debemos trabajar con un tipo particular de filtros, denominados *ultrafiltros*. Decimos que $\mathcal{U}$ es un ultrafiltro sobre un conjunto no vacío I si $\mathcal{U}$ es un filtro propio de I y además

- (3) si $A \subseteq I$, o bien $A \in \mathcal{U}$, o bien $I - A \in \mathcal{U}$.

Encontrar un ultrafiltro sobre un conjunto dado no es una tarea tan sencilla como encontrar simplemente un filtro. Para probar la existencia de un ultrafiltro se requiere utilizar el Lema de Zorn. En particular nos interesará poder encontrar un *ultrafiltro no principal*² sobre $\mathbb{N}$. La demostración de que esto siempre es posible no entra en los objetivos de la monografía, pero se puede encontrar en la Sección 2.6 del libro de Robert Goldblatt [4] para quien esté interesado (allí se demuestra algo más fuerte, que siempre existe un ultrafiltro no principal sobre cualquier conjunto infinito).

Con nuestro ultrafiltro y nuestro concepto del conjunto de coincidencias ya estamos preparados para construir el conjunto de los números hiperreales.

2.3. Manos a la obra

Partimos del conjunto de las sucesiones de números reales, denotado por $\mathbb{R}^{\mathbb{N}}$ y con elementos $r = \langle r_1, r_2, \dots \rangle = \langle r_n : n \in \mathbb{N} \rangle$, o simplemente $r = \langle r_n \rangle$, junto con la suma y el producto usual punto a punto,

$$\langle r_n \rangle + \langle s_n \rangle = \langle r_n + s_n \rangle, \quad \langle r_n \rangle \cdot \langle s_n \rangle = \langle r_n \cdot s_n \rangle.$$

²Decimos que un ultrafiltro es *principal* si contiene un conjunto con un solo elemento.

Tomaremos ahora un ultrafiltro no principal $\mathcal{U}$ sobre $\mathbb{N}$ y definiremos una relación sobre $\mathbb{R}^{\mathbb{N}}$ de la forma

$$r \equiv s \quad \text{si y solo si} \quad \llbracket r = s \rrbracket \in \mathcal{U}.$$

Es sencillo comprobar que esta relación es reflexiva ($r \equiv r$ para toda sucesión $r \in \mathbb{R}^{\mathbb{N}}$) y simétrica (si $r \equiv s$ entonces $s \equiv r$ para todo par de sucesiones $r, s \in \mathbb{R}^{\mathbb{N}}$). Además, también resulta ser transitiva³ (si $r \equiv s$ y $s \equiv t$ entonces $r \equiv t$ para cualesquiera $r, s, t \in \mathbb{R}^{\mathbb{N}}$), por lo cual se tiene que $\equiv$ es una relación de equivalencia, y en virtud de esto podemos agrupar a las sucesiones que están relacionadas entre sí en conjuntos denominados *clases de equivalencia*. Denotamos a la clase de equivalencia de la sucesión r como el conjunto $[r] := \{s \in \mathbb{R}^{\mathbb{N}} : s \equiv r\}$ y al conjunto formado por las clases de equivalencias como ${}^*\mathbb{R} := \{[r] : r \in \mathbb{R}^{\mathbb{N}}\}$. De esta manera, las sucesiones que coincidan en una cantidad grandes de términos estarán representadas por un mismo objeto matemático: su clase de equivalencia.

Estas clases de equivalencias serán lo que llamaremos *números hiperreales*, y el conjunto ${}^*\mathbb{R}$ será el conjunto de los números hiperreales. Pero no nos basta solo con esto, nos gustaría poder realizar operaciones con estos nuevos números, e incluso ordenarlos. Para ello definimos la suma y el producto de las clases de equivalencia como

$$[r] + [s] = [r + s] \quad \text{y} \quad [r] \cdot [s] = [r \cdot s].$$

Un lector avisado se estará preguntado si estas operaciones resultan bien definidas, y la respuesta es que sí lo están. Para comprobarlo basta con ver que $\llbracket r = r' \rrbracket \cap \llbracket s = s' \rrbracket \subseteq \llbracket r + s = r' + s' \rrbracket$ y luego usar las propiedades del ultrafiltro (de manera análoga ocurre con el producto). Respecto al orden, diremos que

$$[r] < [s] \quad \text{si y solo si} \quad \llbracket r < s \rrbracket \in \mathcal{U}.$$

Con todo esto se puede demostrar que el conjunto ${}^*\mathbb{R}$ junto con las operaciones y la relación de orden definidas previamente forma un cuerpo ordenado, con neutro aditivo la clase de equivalencia de la sucesión constantemente cero $[\langle 0, 0, \dots \rangle]$, y neutro multiplicativo la clase de equivalencia de la sucesión constantemente uno $[\langle 1, 1, \dots \rangle]$.

Cabe recalcar que el hecho de que este conjunto sea un cuerpo es muy importante. Eventualmente, cuando utilicemos los números hiperreales para definir la diferenciación, necesitaremos poder dividir por números que representen infinitesimales. Si este conjunto no fuera un cuerpo no habría nada que nos asegurase que el inverso multiplicativo de un infinitesimal exista.

Para aquellos interesados en el álgebra detrás de esta construcción vale la pena aclarar que los números hiperreales no son más que el cociente del anillo de las sucesiones de números reales por la clase de equivalencia $\equiv$, construida en base a un ultrafiltro $\mathcal{U}$. En particular, esta forma de construir a ${}^*\mathbb{R}$ puede encontrarse como “construcción por ultrapotencia” (“ultrapower construction” en inglés). Esto se debe a que el cociente de un producto directo de una estructura algebraica (en este caso, el anillo $\mathbb{R}^{\mathbb{N}}$) mediante una relación de equivalencia generada por un ultrafiltro se denomina de esta manera.

³Basta con verificar que $\llbracket r = s \rrbracket \cap \llbracket s = t \rrbracket \subseteq \llbracket r = t \rrbracket$ y utilizar las propiedades (1) y (2) del filtro.

2.4. ¿Nuevos elementos?

Listo, ya construimos nuestro conjunto ${}^*\mathbb{R}$ y lo dotamos de operaciones y un orden. Ahora enfocaremos nuestros esfuerzos en intentar reconocer cómo se comportan los distintos números hiperreales.

Primeramente nos gustaría poder reconocer a los números reales como un tipo particular de números hiperreales. La solución más fácil es pensar en una función que a cada número real $r \in \mathbb{R}$ lo envíe a la clase de equivalencia formada por la sucesión constantemente r , que denotaremos por simplicidad como $\mathbf{r} := \langle r, r, \dots \rangle$. Es decir, $r \mapsto [\mathbf{r}]$. Por conveniencia denotaremos este mapeo utilizando un asterisco como ${}^*r = [\mathbf{r}]$ (esta notación será recurrente a lo largo del análisis no estándar). Se puede probar que esto es, en efecto, un morfismo inyectivo de $\mathbb{R}$ a ${}^*\mathbb{R}$ que preserva el orden. De esta manera podemos pensar en todos los hiperreales que son de la forma *r para algún $r \in \mathbb{R}$ como una copia de los números reales dentro de los hiperreales. Para simplificar la escritura, llamaremos reales o estándar a los números de la forma *r . Además, en el futuro se evitará frecuentemente la notación *r y simplemente se denotará como r al número, entendiéndose de qué se está hablando por el contexto.

Ahora bien, la pregunta que nos compete y por la cual realizamos toda esta construcción:

¿existen números hiperreales que no sean reales?

La respuesta nuevamente es sí. Veamos dos ejemplos.

Consideremos primero el número hiperreal $\varepsilon = [s]$, donde s es la sucesión

$$s = \left\langle 1, \frac{1}{2}, \frac{1}{3}, \dots \right\rangle = \left\langle \frac{1}{n} \right\rangle.$$

Es claro que este número resulta positivo, ${}^*0 < \varepsilon$, ya que $\llbracket 0 < s \rrbracket = \mathbb{N} \in \mathcal{U}$. Tomemos ahora un número real positivo cualquiera *r . Por la propiedad arquimediana es fácil ver que el conjunto $\llbracket s < \mathbf{r} \rrbracket = \{n \in \mathbb{N} : \frac{1}{n} < r\}$ es cofinito. Aquí es donde entra en juego que el ultrafiltro $\mathcal{U}$ que elegimos no es principal. Se puede demostrar que si $\mathcal{U}$ es no principal entonces contiene a todos los conjuntos cofinitos. De esta manera $\llbracket s < \mathbf{r} \rrbracket \in \mathcal{U}$ y se sigue que $\varepsilon < {}^*r$ para todo $r \in \mathbb{R}$. Por lo tanto, podemos concluir que ε es un número hiperreal que no es un número real. Además, es notable ver como ε encaja con la descripción de número infinitesimal que deseábamos, pues es positivo y menor que cualquier número real.

El otro ejemplo que consideraremos será $\omega = [t]$, con la sucesión

$$t = \langle 1, 2, 3, \dots \rangle = \langle n \rangle.$$

Usando la propiedad arquimediana se puede ver que $\llbracket \mathbf{r} < t \rrbracket = \{n \in \mathbb{N} : r < n\}$ es cofinito para todo $r \in \mathbb{R}$. De esta manera, se tiene que ω es mayor que cualquier número real, y en consecuencia no puede ser un número real.

Es importante destacar que la no principalidad del ultrafiltro escogido es clave para que existan nuevos elementos a la hora de construir los números reales. Se puede probar que si se escogiera un ultrafiltro principal entonces la ultrapotencia ${}^*\mathbb{R}$ sería isomorfa a $\mathbb{R}$, y por lo tanto no tendríamos nuevos elementos como ε o ω .

Si bien la definición de entidades matemáticas que involucren números hiperreales (como conjuntos, funciones o sucesiones) no presenta ninguna dificultad, nos resultará más útil

estudiar un tipo particular de objetos, aquellos que se definan en base a números reales y puedan ser *extendidos* a números hiperreales.

Pensemos en un conjunto de números reales $A \subseteq \mathbb{R}$. Visto que los números reales se consideran incluidos en los números hiperreales, entonces también resulta que $A \subseteq {}^*\mathbb{R}$. Pero iremos más allá. Definiremos otro subconjunto *A de ${}^*\mathbb{R}$ de la siguiente forma

$$[r] \in {}^*A \quad \text{si y solo si} \quad \llbracket r \in A \rrbracket := \{n \in \mathbb{N} : r_n \in A\} \in \mathcal{U}.$$

Llamaremos al conjunto *A como el *conjunto extendido* o la *extensión* de A .

Según esta definición es claro que $A \subseteq {}^*A$. Los elementos de A (que resultan ser todos números reales) son los que llamaremos *elementos estándar* del conjunto *A , y de manera similar llamaremos *elementos no estándar* a los que pertenecen a ${}^*A - A$. Para ilustrar esto tomemos como ejemplo el conjunto de los números naturales $\mathbb{N}$ pensado como subconjunto de $\mathbb{R}$. Ya vimos que el número hiperreal $\omega = \langle 1, 2, \dots \rangle$ no es real, por lo cual $\omega \notin \mathbb{N}$, pero en cambio $\omega \in {}^*\mathbb{N}$ ya que todos los elementos de la sucesión $\langle 1, 2, \dots \rangle$ son naturales. De esta forma, ω es un elemento no estándar de ${}^*\mathbb{N}$.

No todos los conjuntos extendidos tienen elementos no estándar. En ocasiones el conjunto ${}^*A - A$ puede ser vacío. Esto ocurre cuando el conjunto A es finito. Más aún, esta es una caracterización, lo que significa *A tiene términos no estándar si y solo si A es infinito.

Pensemos ahora en una función $f : \mathbb{R} \rightarrow \mathbb{R}$. La manera de extenderla a otra función ${}^*f : {}^*\mathbb{R} \rightarrow {}^*\mathbb{R}$ es definirla de la forma

$${}^*f(\llbracket \langle r_n \rangle \rrbracket) := \llbracket \langle f(r_n) \rangle \rrbracket.$$

A simple vista la notación puede resultar enrevesada, pero es bastante sencillo de verlo con un ejemplo. Tomemos la función $f(x) = x^2$, entonces

$${}^*f(\omega) = {}^*f(\llbracket \langle 1, 2, 3, \dots \rangle \rrbracket) = \llbracket \langle 1, 4, 9, \dots \rangle \rrbracket.$$

Si en lugar de tener una función con dominio $\mathbb{R}$ tenemos $f : A \rightarrow \mathbb{R}$, con $A \subseteq \mathbb{R}$. Extenderemos la función a ${}^*f : {}^*A \rightarrow {}^*\mathbb{R}$ de la siguiente forma: Para $[r] \in {}^*A$ definiremos

$$s_n = \begin{cases} f(r_n), & \text{si } n \in \llbracket r \in A \rrbracket, \\ 0, & \text{si } n \notin \llbracket r \in A \rrbracket. \end{cases}$$

Así, la función extendida será ${}^*f([r]) = [s]$, o bien ${}^*f(\llbracket \langle r_n \rangle \rrbracket) = \llbracket \langle f(r_n)\chi_A(r_n) \rangle \rrbracket$ ⁴. Notemos que la forma de definir la sucesión cuando $n \notin \llbracket r \in A \rrbracket$ (es decir cuando $r_n \notin A$) es totalmente arbitraria. Todas las posibles formas de definir la sucesiones diferirán entre sí en un conjunto que no pertenece al ultrafiltro, es decir un conjunto que no es grande, y en consecuencia su clase de equivalencia será la misma, representando así al mismo número hiperreal.

Tomemos como ejemplo la función $f : (0, \infty) \rightarrow \mathbb{R}$ definida como $f(x) = x^{-1}$. Su extensión *f estará definida sobre el conjunto ${}^*(0, \infty)$. Si tomamos la sucesión

$$t = \langle -1, 1, 3, 5, 7, \dots \rangle = \langle 2n - 3 \rangle,$$

⁴Aquí χ_A es la función característica del conjunto A . Es decir $\chi_A(x) = 1$ si $x \in A$ y $\chi_A(x) = 0$ si $x \notin A$.

entonces el hiperreal $[t] \in {}^*(0, \infty)$ tendrá por imagen

$${}^*f([t]) = \left[\left\langle 0, 1, \frac{1}{3}, \frac{1}{5}, \frac{1}{7}, \dots \right\rangle \right].$$

Por último, si tenemos una función real de varias variables $f : A_1 \times \dots \times A_k \rightarrow \mathbb{R}$ con $A_1, \dots, A_k \subseteq \mathbb{R}$, su extensión será ${}^*f : {}^*A_1 \times \dots \times {}^*A_k \rightarrow {}^*\mathbb{R}$, y de la misma forma que antes para $[r^{(1)}] \in {}^*A_1, \dots, [r^{(k)}] \in {}^*A_k$ definiremos

$$s_n = \begin{cases} f(r_n^{(1)}, \dots, r_n^{(k)}), & \text{si } n \in \llbracket r^{(1)} \in A_1 \rrbracket \cap \dots \cap \llbracket r^{(k)} \in A_k \rrbracket, \\ 0, & \text{si } n \notin \llbracket r^{(1)} \in A_1 \rrbracket \cap \dots \cap \llbracket r^{(k)} \in A_k \rrbracket. \end{cases}$$

Luego, definimos ${}^*f([r^{(1)}], \dots, [r^{(k)}]) = [s]$.

Si bien esta definición puede resultar difícil de procesar a simple vista, es tan simple como pensar en que las funciones al extenderse actúan punto a punto en los elementos de la sucesión. Por ejemplo, ya hemos definido cómo sumar y cómo multiplicar dos hiperreales, pero si ahora quisiéramos definir a^b con $a, b \in {}^*\mathbb{R}$ y $a > 0$, bastaría con interpretar esto como una función de dos variables $f : \mathbb{R}^+ \times \mathbb{R} \rightarrow \mathbb{R}$ definida como $f(x, y) = x^y$, y extenderla como vimos previamente. Por ejemplo, si tomamos

$$\omega = [\langle 1, 2, 3, 4, \dots \rangle], \quad \xi = [\langle 1, -1, 1, -1, \dots \rangle],$$

se tiene que

$$\omega^\xi = [\langle 1^1, 2^{-1}, 3^1, 4^{-1}, \dots \rangle] = \left[\left\langle 1, \frac{1}{2}, 3, \frac{1}{4}, \dots \right\rangle \right].$$

Además, podemos notar que $\xi = -(-1)^\omega$.

En lo que resta de la monografía omitiremos la notación *f a la hora de hablar de la extensión de una función f . Se entenderá por el contexto y por el conjunto al que pertenece el número en que se evalúa, si se trata de la función real o su extensión hiperreal.

2.5. Clasificación y propiedades

Introduciremos ahora un par de definiciones que nos permitirán clasificar a los números hiperreales.

Diremos que un número hiperreal x es:

- *limitado* si existen dos números reales r y s tales que $r < x < s$.
- *ilimitado positivo* si $x > r$ para todo real r .
- *ilimitado negativo* si $x < r$ para todo real r .
- *ilimitado* si es ilimitado positivo o ilimitado negativo.
- *infinitesimal positivo* si $0 < x < r$ para todo real positivo r .
- *infinitesimal negativo* si $r < x < 0$ para todo real negativo r .

- *infinitesimal* si es infinitesimal positivo, infinitesimal negativo o 0.
- *apreciable* si es limitado pero no infinitesimal.

Una observación importante es que, a partir de ahora, el concepto de infinitesimal que adoptaremos es más amplio que el que utilizamos en el razonamiento intuitivo al inicio de esta monografía. Aquella primera noción se corresponde con lo que ahora llamamos *infinitesimal positivo*. Podemos reformular la definición inicial diciendo que “un infinitesimal es una cantidad cuyo valor absoluto es menor que cualquier número real positivo”. Con esta nueva perspectiva, el concepto abarca tanto al cero como a los infinitesimales negativos.

Denotaremos al conjunto de los números hiperreales limitados como $\mathbb{L}$, y al conjunto de los infinitesimales como $\mathbb{I}$. Además, se denotará con A_∞ al conjunto de los hiperreales ilimitados de un conjunto dado A .

El último concepto que introduciremos en este capítulo es uno extremadamente útil. Se trata de la *sombra* o la *parte estándar*. Pero antes de entrar en el tema veamos una forma de relacionar los números hiperreales que resulta bastante intuitiva.

Diremos que un número hiperreal x se encuentra *infinitesimalmente cerca* de otro hiperreal y si $x - y$ es infinitesimal. Lo denotamos por $x \simeq y$. Resulta ser que $\simeq$ es una relación de equivalencia, y a las clases de equivalencia las denominamos *halos*, de manera que el halo de un hiperreal x se define como $\text{hal}(x) := \{y \in {}^*\mathbb{R} : x \simeq y\}$. Es importante a notar que si x e y son ambos reales, entonces no se encuentran infinitesimalmente cerca, ya que su diferencia es real.

Existe un teorema fundamental en análisis no estándar que afirma que todo número hiperreal limitado se encuentra en el halo de un número real. Más aún, este número real es único para cada hiperreal, debido a que los halos de números reales son disjuntos entre sí. De esta manera, dado un número hiperreal limitado $x \in \mathbb{L}$, llamaremos *sombra* de x a este real infinitesimalmente cerca de x , y lo denotaremos por $\text{sh}(x)$ (por *shadow* en inglés). Es común en varios textos de análisis no estándar llamar a la sombra de x como la *parte estándar* de x , denotándolo como $\text{st}(x)$. A lo largo de la monografía utilizaremos indistintamente ambas formas de nombrarlo, según se crea conveniente, pero para evitar confusiones siempre la denotaremos por $\text{sh}(\cdot)$. La demostración de este teorema no resulta muy complicada. Se basa en definir el conjunto $A = \{r \in \mathbb{R} : r < x\}$ para $x \in \mathbb{L}$, probar que es no vacío y acotado superiormente, y utilizar la completitud de los números reales para demostrar que la sombra de x resulta ser el supremo de A . De hecho, se puede probar que la existencia de la sombra no solamente es consecuencia de la completitud de $\mathbb{R}$, sino que es equivalente a la misma.

El concepto de halo es muy curioso. Así por como está definido se ve bastante sencillo, pero intente el lector imaginarse como luciría un halo en la recta numérica. Uno tiende a imaginárselo como un intervalo abierto centrado en un número, pero muy pequeño, demasiado, tanto que no se puede ni siquiera dibujar ni representar gráficamente. Ahora intente imaginarse el hecho de que cada número real tiene un halo centrado en este, en el cual hay una cantidad enorme, infinita, de números hiperreales, y que todos esos halos son disjuntos entre sí. Casi parecería como que los números reales fueran un conjunto discreto, en el que podríamos separar todos sus elementos entre sí. Nada más alejado de la realidad.

Para concluir esta sección, presentaremos algunas propiedades aritméticas de los números hiperreales que serán de uso frecuente a lo largo de la monografía. Sean δ, μ infinitesimales; a, b apreciables; x, y, z limitados y H, K ilimitados. Entonces:

- (1) $x + y$ y $x \cdot y$ son limitados.⁵
- (2) $\delta + \mu$ y $\delta \cdot x$ son infinitesimales.⁶
- (3) $H + x$ son ilimitados.
- (4) $\frac{\delta}{a}$, $\frac{\delta}{H}$ y $\frac{x}{H}$ son infinitesimales.
- (5) $\frac{x}{a}$ es limitado, y en particular $\frac{b}{a}$ es apreciable.
- (6) $\frac{a}{\delta}$, $\frac{H}{\delta}$ y $\frac{H}{a}$ son ilimitados (siempre que $\delta \neq 0$).
- (7) Como 1 es apreciable, se tiene que H^{-1} es infinitesimal, a^{-1} es apreciable y δ^{-1} es ilimitado (si $\delta \neq 0$).
- (8) Si $x \leq y$, entonces $\text{sh}(x) \leq \text{sh}(y)$.
- (9) Si $x \simeq y$, entonces $xz \simeq yz$.

Si bien no demostraremos estas propiedades, resulta una tarea recomendable para el lector detenerse en ellas y reflexionar sobre lo que significan. Algunas quizás luzcan difíciles de comprender a simple vista, y otras, muy probablemente, rememoren distintas “reglas” típicas del cálculo estándar. Por ejemplo, decir que “ $\frac{1}{\infty} = 0$ ”, se puede relacionar con la propiedad (7), interpretando como “ ∞ ” a un número ilimitado positivo y como “0” un número infinitesimal. De manera similar, la propiedad (3) tiene su análogo estándar como “ $\infty + x = \infty$ ” para cualquier x real.

Por último, presentaremos algunas formas indeterminadas que pueden surgir al operar con números hiperreales:

$$\frac{\delta}{\mu}, \quad \frac{H}{K}, \quad \delta \cdot H, \quad H - K.$$

Siguiendo la analogía con las “reglas” del cálculo estándar, resulta evidente la relación con las formas indeterminadas ya conocidas:

$$\frac{0}{0}, \quad \frac{\infty}{\infty}, \quad 0 \cdot \infty, \quad \infty - \infty.$$

Analicemos, por ejemplo, la última operación. Si tomamos M hiperreal ilimitado, entonces se tiene que $M - M = 0$, que es infinitesimal. Sin embargo, como $2M$ también es ilimitado, ocurre que $2M - M = M$, que sigue siendo ilimitado. Esto muestra que la resta de dos hiperreales ilimitados puede dar como resultado tanto un número infinitesimal como uno también ilimitado. Encontrar ejemplos similares para las demás formas indeterminadas no es una tarea difícil, y se invita al lector a intentarlo si así lo desea.

Es de destacar que, en el marco del análisis no estándar, este tipo de ambigüedad se hace especialmente evidente: las indeterminaciones no requieren desarrollos límite para ser comprendidas, sino que surgen de manera directa al operar con los objetos mismos.

⁵Por lo que $(\mathbb{L}, +, \cdot)$ es un anillo.

⁶De esta manera, $\mathbb{I}$ resulta ser un ideal de $\mathbb{L}$.

2.6. El principio de transferencia

En lo que sigue de esta monografía vamos a utilizar fuertemente un resultado de lógica que nos facilitará muchas pruebas y deducciones. Se trata del *principio de transferencia*.

Coloquialmente, el principio de transferencia establece que todas las afirmaciones matemáticas verdaderas sobre los números reales que estén formuladas en lenguaje de primer orden son también válidas para los números hiperreales, y viceversa.

Desde un punto de vista un poco más formal, lo que nos va a permitir el principio de transferencia es decir que una fórmula bien formada de primer orden φ sobre los números reales tiene el mismo valor de verdad que la fórmula bien formada de primer orden $^*\varphi$ sobre los números hiperreales. Esta última resulta ser la *extensión* de φ .

La extensión $^*\varphi$ de la fórmula φ es simplemente extender (valga la redundancia) todos los objetos matemáticos que se involucren (como pueden ser conjuntos, funciones, relaciones, proposiciones, etc). Por ejemplo, si

$$\varphi = (\forall x \in \mathbb{R}^+)(\forall y \in \mathbb{R}^+)(\log(xy) = \log(x) + \log(y)),$$

entonces su extensión será

$$^*\varphi = (\forall x \in {}^*\mathbb{R}^+)(\forall y \in {}^*\mathbb{R}^+)(^*\log(xy) = ^*\log(x) + ^*\log(y)).$$

Usualmente se suele omitir el símbolo $*$ en la extensión del log, pero aquí se incluyó con fines demostrativos.

Vale la pena destacar que el principio de transferencia también nos permite establecer equivalencias entre fórmulas del tipo

$$(\forall x \in A)\psi(x) \quad \longleftrightarrow \quad (\forall x \in A) {}^*\psi(x)$$

y

$$(\exists x \in A)\psi(x) \quad \longleftrightarrow \quad (\exists x \in A) {}^*\psi(x).$$

Es muy importante tener en cuenta que “la parte que se transforma” de estándar a no estándar debe ser la de la derecha, pues sería incorrecto transformar el conjunto en que se aplican los cuantificadores sin transformar la fórmula en la que se encuentra la variable sobre la que actúan los mismos. A continuación veremos dos ejemplos de por qué esto no se cumple.

- Es cierto que $(\forall r \in \mathbb{R})(\exists N \in \mathbb{N})(r < N)$, ya que es la propiedad arquimediana, pero es falso que $(\forall r \in {}^*\mathbb{R})(\exists N \in \mathbb{N})(r < N)$ (basta con elegir un hiperreal r ilimitado positivo). Lo correcto sería afirmar que $(\forall r \in {}^*\mathbb{R})(\exists N \in {}^*\mathbb{N})(r < N)$.
- Es falso que $(\exists b \in \mathbb{R})(\forall x \in \mathbb{R})(x < b)$, pero es cierto que $(\exists b \in {}^*\mathbb{R})(\forall x \in \mathbb{R})(x < b)$ (basta con elegir un hiperreal b ilimitado positivo).

La demostración de por qué se cumple el principio de transferencia involucra conceptos de lógica de primer orden en los cuales no entraremos en detalle. El lector interesado puede recurrir al capítulo 4 de [4] para encontrar un desarrollo con mayor profundidad.

3. Análisis no estándar

“Make-believe it’s hyper real.”

— Lorde, *Buzzcut Season*

Ahora que ya definimos los números hiperreales y sabemos cómo operar con ellos, los emplearemos para reproducir algunos resultados clásicos del análisis. Nuestro método de trabajo será el siguiente: primero encontraremos equivalencias entre definiciones del análisis *estándar* y ciertas propiedades de nuestros objetos *no estándar* en base al principio de transferencia. Luego, demostraremos los teoremas habituales empleando las entidades no estándar que hemos construido. Por ejemplo, si demostramos que la continuidad de una función real equivale a que su extensión no estándar cumpla una propiedad P , y que su diferenciabilidad equivale a otra propiedad Q , entonces bastará probar (en un contexto no estándar) que Q implica P . De este modo obtendremos el teorema clásico: toda función diferenciable es continua. En algunos casos la demostración *no estándar* puede resultar más sencilla que la *estándar*.

3.1. Sucesiones

Comenzaremos con uno de los objetos matemáticos con los que primero se trabaja en un curso de Análisis: las sucesiones.

Podemos representar a las sucesiones de números reales $\langle s_n \rangle$ simplemente como una función $s : \mathbb{N} \rightarrow \mathbb{R}$, de forma que $s(n) = s_n$. Así, es fácil pensar en su extensión a una *hipersucesión* como la extensión de la función $s : {}^*\mathbb{N} \rightarrow {}^*\mathbb{R}$. De esta manera, tendrá sentido pensar en s_n , aún cuando n no sea natural.

Los números pertenecientes al conjunto ${}^*\mathbb{N}$ se denominan *hipernaturales*. Se puede probar que los únicos elementos no estándar de ${}^*\mathbb{N}$ son ilimitados (y en particular ilimitados positivos), es decir ${}^*\mathbb{N} - \mathbb{N} = {}^*\mathbb{N}_\infty$. Si $n \in {}^*\mathbb{N}_\infty$ entonces el término de la hipersucesión s_n se llama *término extendido*, y el conjunto $\{s_n : n \in {}^*\mathbb{N}_\infty\}$ se denomina *cola extendida*.

Veamos como lucen los términos extendidos de una sucesión. Para ello analicemos la sucesión

$$s = \left\langle 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots \right\rangle,$$

y los números hipernaturales

$$\omega = [1, 2, 3, 4, \dots], \quad P = 2^\omega = [2, 4, 8, 16, \dots], \quad Q = \frac{P}{4} = \left[\left\langle \frac{1}{2}, 1, 2, 4, \dots \right\rangle \right].$$

De esta manera, los términos extendidos de la sucesión correspondiente a los distintos números hipernaturales ilimitados son

$$s_\omega = \left[\left\langle 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots \right\rangle \right], \quad s_P = \left[\left\langle \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \dots \right\rangle \right], \quad s_Q = \left[\left\langle 0, 1, \frac{1}{2}, \frac{1}{4}, \dots \right\rangle \right].$$

Podemos notar que el término extendido de una hipersucesión es el hiperreal dado por la clase de equivalencia de otra sucesión formada por un cantidad grande de los términos

estándar de la sucesión original, pero reordenados⁷. En particular es curioso ver como siempre s_ω resulta ser la clase de equivalencia de la misma sucesión s . A partir de ahora nos reservaremos ω para denotar siempre al hiperreal $\omega = [1, 2, 3, \dots]$.

Ya estamos en condiciones de enunciar nuestro primer teorema importante.

Teorema 3.1. *Una sucesión de números reales $s = \langle s_n \rangle$ converge a $L \in \mathbb{R}$ si y solo si $s_n \simeq L$ para todo $n \in {}^*\mathbb{N}_\infty$.*

Bosquejo de la demostración.

Si la sucesión s converge a L entonces para todo $\varepsilon \in \mathbb{R}^+$ existe $N \in \mathbb{N}$ tal que para todo $n \in \mathbb{N}$ se tiene que si $n \geq N$ entonces $|s_n - L| < \varepsilon$. Escrito simbólicamente:

$$(\forall \varepsilon \in \mathbb{R}^+)(\exists N \in \mathbb{N})(\forall n \in \mathbb{N})(n \geq N \implies |s_n - L| < \varepsilon).$$

Luego, por el principio de transferencia podemos afirmar que esto mismo sucede para todo $n \in {}^*\mathbb{N}$, es decir

$$(\forall \varepsilon \in \mathbb{R}^+)(\exists N \in \mathbb{N})(\forall n \in {}^*\mathbb{N})(n \geq N \implies |s_n - L| < \varepsilon).^8$$

En particular, si $n \in {}^*\mathbb{N}_\infty$ implica que $n \geq N$ para cualquier $N \in \mathbb{N}$. De esta manera se tiene que para cualquier $\varepsilon \in \mathbb{R}^+$ resulta que $|s_n - L| < \varepsilon$, y por lo tanto $s_n \simeq L$.

Supongamos ahora que $s_n \simeq L$ para todo $n \in {}^*\mathbb{N}_\infty$. Tomemos $\varepsilon \in \mathbb{R}^+$ arbitrario y fijemos un hipernatural ilimitado $N \in {}^*\mathbb{N}_\infty$. Está claro que para cualquier $n \in {}^*\mathbb{N}$ sucede que si $n \geq N$ entonces $|s_n - L| < \varepsilon$, y por lo tanto

$$(\forall \varepsilon \in \mathbb{R}^+)(\exists N \in {}^*\mathbb{N})(\forall n \in {}^*\mathbb{N})(n \geq N \implies |s_n - L| < \varepsilon).$$

Por el principio de transferencia se sigue que también existe $N \in \mathbb{N}$ que cumple con lo anterior. Así, se obtiene

$$(\forall \varepsilon \in \mathbb{R}^+)(\exists N \in \mathbb{N})(\forall n \in \mathbb{N})(n \geq N \implies |s_n - L| < \varepsilon),$$

que resulta en la definición de que la sucesión s converja a L . ■

Este resultado es muy importante, pues podremos probar que una sucesión es convergente sin necesidad de hacer un argumento de límite, simplemente demostrando que todos los términos extendidos tienen la misma parte estándar. Obviamente esto último puede resultar en ocasiones mucho más difícil que hacer una simple demostración estándar, pero a veces es más sencillo.

A continuación veremos que si $c \in (0, 1)$, entonces la sucesión $s = \langle c^n \rangle$ converge a 0. Para ello primero utilizaremos un resultado que afirma que, si tenemos una sucesión

⁷Si alguno de los términos de la sucesión que define al hipernatural resulta no ser natural, la sucesión que define al término extendido asociado a este hipernatural tendrá un 0 en esa misma posición. Esto se puede apreciar en el caso de s_Q , debido a que el primer término de la sucesión que define a Q es $\frac{1}{2}$. Notemos que este hecho no resulta ningún inconveniente, pues un hipernatural tiene, a lo sumo, una cantidad no grande de términos no naturales.

⁸La proposición $(n \geq N \implies |s_n - L| < \varepsilon)$ en esta expresión debe interpretarse como la versión extendida de la misma. A pesar de que no haya ningún símbolo $*$ que lo denote explícitamente, se sobreentiende debido a que $n \in {}^*\mathbb{N}$.

monótona decreciente acotada inferiormente, entonces esta sucesión converge. En particular, se tiene que converge al ínfimo del conjunto formado por los términos de la sucesión. Esta afirmación puede ser probada utilizando argumentos no estándar, pero para no extender la monografía se decidió no entrar en detalle⁹. Como s es monótona decreciente y está acotada inferiormente por 0, entonces podemos afirmar que convergerá a un número real $L = \inf\{c^n : n \in \mathbb{N}\}$. Lo que demostraremos será que este ínfimo resulta ser 0. Para ello notemos que, si $N \in {}^*\mathbb{N}_\infty$ sucede que $c^N \simeq L$, y así $c^N c \simeq Lc$. Luego, como $c^{N+1} = c^N c$ ¹⁰, al ser $N+1$ también un hipernatural ilimitado se sigue que $c^N c \simeq L$. De esta manera, $Lc \simeq L$, y como tanto L como Lc son reales, se tiene la igualdad $Lc = L$. Finalmente, como $c \neq 1$ solo resta que $L = 0$.

Otro uso notable del teorema anterior es obtener el resultado de ciertas operaciones con números hiperreales. Por ejemplo, es bien sabido que la sucesión

$$s = \left\langle \left(1 + \frac{1}{n}\right)^n \right\rangle$$

converge al número e (el número de Euler). Por lo cual $s_N \simeq e$ para todo $N \in {}^*\mathbb{N}_\infty$, en particular $\text{sh}(s_\omega) = e$. Por como definimos las operaciones entre números hiperreales tenemos que

$$\left(1 + \frac{1}{\omega}\right)^\omega = \left[\left\langle \left(1 + \frac{1}{n}\right)^n \right\rangle \right] = [s] = s_\omega.$$

Así, podemos afirmar $(1 + 1/\omega)^\omega \simeq e$, o equivalentemente $\text{sh}((1 + 1/\omega)^\omega) = e$. Si ahora definimos $\varepsilon = \omega^{-1}$, una manera análoga de escribir este resultado es

$$(1 + \varepsilon)^{\frac{1}{\varepsilon}} \simeq e.$$

Resulta oportuno también mencionar otros dos teoremas importantes, el primero sobre puntos de acumulación y el segundo sobre sucesiones acotadas.

Teorema 3.2. *$L \in \mathbb{R}$ es un punto de acumulación de la sucesión $s = \langle s_n \rangle$ si y solo si existe algún $N \in {}^*\mathbb{N}_\infty$ tal que $s_N \simeq L$.*

La demostración de este teorema es bastante similar a la del Teorema 3.1, utilizando la definición de punto de acumulación en lugar de convergencia. Este teorema también saca a la luz una comparación muy curiosa: decir que L es un punto de acumulación de s es mucho más débil que decir que s converge a L , pero equivale a decir que existe una subsucesión de s que converge a L . De la misma forma, pedir que exista al menos un término extendido que esté infinitesimalmente cerca de L es mucho más débil que pedir que todos los términos extendidos lo estén. Si $\langle s_{n_k} \rangle_{k \in \mathbb{N}}$ es una subsucesión de s que converge a L , entonces si pensamos en el número hipernatural que resulta de la clase de equivalencia del subconjunto de índices de esta subsucesión $N = [n_k]_{k \in \mathbb{N}}$ se tiene que $s_N = [s_{n_k}]_{k \in \mathbb{N}} \simeq L$. De esta forma podemos completar la intuición que generamos al principio de esta sección y concluir que cada término extendido de nuestra hipersucesión

⁹El decrecimiento monótono de la sucesión equivale a que los términos extendidos no sean ilimitados positivos, y el hecho de que la sucesión esté acotada inferiormente se traduce en que los mismo tampoco sean ilimitados negativos. Ambas propiedades combinadas permiten probar que la cola extendida se encuentra infinitesimalmente cerca de la máxima cota inferior de la sucesión, es decir el ínfimo.

¹⁰Esta propiedad resulta ser la extensión por principio de transferencia de que para todo $n \in \mathbb{N}$ se tiene que $c^{n+1} = c^n c$

está asociado al comportamiento de una subsucesión correspondiente, de manera que la convergencia de la subsucesión se traduce en que el término extendido sea limitado.

Veamos un último teorema sobre sucesiones:

Teorema 3.3. *Una sucesión de números reales $s = \langle s_n \rangle$ está acotada si y solo si todos sus términos extendidos son limitados.*

La idea detrás de la demostración del Teorema 3.3 es ver que la acotación de todos los términos de la sucesión por un número real es equivalente a que todos los términos de la hipersucesión (entre ellos los extendidos) estén acotados por el mismo número real.

Ya hemos visto que todo número limitado está infinitesimalmente cerca de un número real, su sombra. Por lo tanto, es inmediato ver que, si todos los términos extendidos de una hipersucesión son limitados, entonces la sombra de cada uno de ellos es un punto de acumulación, por el Teorema 3.2. De esta manera, casi sin darnos cuenta, hemos demostrado el teorema de Bolzano-Weierstrass, que afirma que toda sucesión de números reales acotada tiene al menos un punto de acumulación, pues al tener una sucesión acotada, por el Teorema 3.3 sus términos extendidos son limitados, y por la deducción hecha previamente existe al menos un punto de acumulación.

3.2. Continuidad

Así como hicimos con las sucesiones, mencionaremos un teorema que nos permitirá caracterizar las funciones reales continuas a través de una propiedad de su extensión.

Teorema 3.4. *Una función $f : A \rightarrow \mathbb{R}$ con $A \subseteq \mathbb{R}$ es continua en $c \in A$ si y solo si para todo $x \in {}^*A$ tal que $x \simeq c$ sucede que $f(x) \simeq f(c)$.*

Bosquejo de la demostración.

Si f es continua en c entonces, por la definición estándar, para todo $\varepsilon \in \mathbb{R}^+$ existe $\delta \in \mathbb{R}^+$ tal que para todo $x \in A$ sucede que si $|x - c| < \delta$ entonces $|f(x) - f(c)| < \varepsilon$. Luego, por transferencia, esto ocurre para todo $x \in {}^*A$, es decir,

$$(\forall \varepsilon \in \mathbb{R}^+)(\exists \delta \in \mathbb{R}^+)(\forall x \in {}^*A)(|x - c| < \delta \implies |f(x) - f(c)| < \varepsilon).$$

Tomando un valor $x \in {}^*A$ y para ε dado, si sucede que $x \simeq c$, entonces $|x - c| < \delta$ cualquiera sea el valor de δ , por lo cual $|f(x) - f(c)| < \varepsilon$. Como esto ocurre para cualquier valor de $\varepsilon \in \mathbb{R}^+$ arbitrario, concluimos que $f(x) \simeq f(c)$ siempre que $x \simeq c$.

Supongamos ahora que para todo $x \in {}^*A$ se tiene que $x \simeq c$ implica $f(x) \simeq f(c)$. Si tomamos $\varepsilon \in \mathbb{R}^+$ arbitrario, para cualquier δ infinitesimal positivo queelijamos ocurre que si $|x - c| < \delta$ entonces $|f(x) - f(c)| < \varepsilon$. De esta manera se tiene que

$$(\forall \varepsilon \in \mathbb{R}^+)(\exists \delta \in {}^*\mathbb{R}^+)(\forall x \in {}^*A)(|x - c| < \delta \implies |f(x) - f(c)| < \varepsilon).$$

Finalmente, aplicando equivalencia se obtiene una fórmula que resulta ser la definición estándar de continuidad. ■

A la hora de demostrar que una función $f : A \rightarrow \mathbb{R}$ es continua en $c \in A$ usando esta caracterización no estándar suele ser conveniente escribir a $x \in {}^*A$ tal que $x \simeq c$ como $x = c + \varepsilon$ con $\varepsilon \simeq 0$ (es decir, ε infinitesimal). Veamos algunos ejemplos.

- Pensemos en la función $f : \mathbb{R} \rightarrow \mathbb{R}$ definida por $f(x) = x^2$. Sean $c \in \mathbb{R}$ y $x \in {}^*\mathbb{R}$ de la forma $x = c + \varepsilon$ con $\varepsilon \simeq 0$. Así, $f(x) = (c + \varepsilon)^2 = c^2 + 2c\varepsilon + \varepsilon^2 \simeq c^2 = f(c)$, pues $2c\varepsilon$ y ε^2 son infinitesimales. De esta manera concluimos que f es continua en todo $c \in \mathbb{R}$.
- Sea $a \in \mathbb{R}^+$, veamos que la función $f : \mathbb{R} \rightarrow \mathbb{R}$ definida como $f(x) = a^x$ es continua en todo $\mathbb{R}$. Para ello veamos primero que, para todo $\varepsilon \simeq 0$ resulta que $a^\varepsilon \simeq 1$ ¹¹. Tomemos $r \in \mathbb{R}^+$, sin pérdida de generalidad $0 < r < 1$. Luego $0 < 1 - r < 1 < 1 + r$ y en consecuencia $\log_a(1 - r) < 0 < \log_a(1 + r)$. Como ε es infinitesimal $\log_a(1 - r) < \varepsilon < \log_a(1 + r)$. Finalmente, $1 - r < a^\varepsilon < 1 + r$, y así $|a^\varepsilon - 1| < r$. De esta manera concluimos que $a^\varepsilon \simeq 1$.

Ahora sí, volvamos a nuestra función $f(x) = a^x$. Sean $c \in \mathbb{R}$ y $x = c + \varepsilon$ con $\varepsilon \simeq 0$. Luego,

$$f(x) = a^x = a^{c+\varepsilon} = a^c a^\varepsilon \simeq a^c \cdot 1 = a^c = f(c).$$

De esta forma $f(x) \simeq f(c)$ y f es continua en todo $c \in \mathbb{R}$.

- Probaremos ahora que la función $f(x) = \sin(x)$ es continua en todo $\mathbb{R}$. Sea $c \in \mathbb{R}$ y $x = c + \varepsilon$ con $\varepsilon \simeq 0$. Por propiedad del seno de la suma¹² se tiene que

$$\sin(x) = \sin(c + \varepsilon) = \sin(c) \cos(\varepsilon) + \cos(c) \sin(\varepsilon).$$

Es posible probar con argumentos no estándar que $\cos(\varepsilon) \simeq 1$ y $\sin(\varepsilon) \simeq 0$ para $\varepsilon \simeq 0$. De esta manera

$$\sin(c) \cos(\varepsilon) + \cos(c) \sin(\varepsilon) \simeq \sin(c) \cdot 1 + \cos(c) \cdot 0 = \sin(c).$$

Así, $f(x) = \sin(c + \varepsilon) \simeq \sin(c) = f(c)$ y f es continua en todo $c \in \mathbb{R}$.

Meditemos un poco sobre qué nos quiere decir esta caracterización de las funciones continuas. Está claro que una interpretación bastante directa es que las funciones reales continuas son aquellas cuya extensión cumple con la propiedad de que todos los números hiperreales “cerca” a números reales se mantienen “cerca” luego de aplicar la función. Dicho simbólicamente, si nuestra función tiene por dominio A , entonces f es continua si y solo si $f(\text{hal}(c)) \subseteq \text{hal}(f(c))$ para todo $c \in A$. Es decir, nuestra función no “separa” elementos dentro del halo de un real.

Una pregunta que surge naturalmente es ¿se sigue cumpliendo esta propiedad en las funciones continuas para halos de números no reales? La respuesta en general es no, esta es una propiedad más fuerte que la caracterización de la continuidad. Consideremos la función

$$f(x) = \text{sgn}(x) = \begin{cases} 1, & \text{si } x > 0, \\ 0, & \text{si } x = 0, \\ -1, & \text{si } x < 0. \end{cases}$$

Esta claro que f es continua en $\mathbb{R}^+$, por lo que se tiene que para todo $c \in \mathbb{R}^+$, $f(\text{hal}(c)) \subseteq \text{hal}(f(c))$. Tomamos ahora el hiperreal $\varepsilon = [1, 1/2, 1/3, \dots]$. Se puede ver que $\varepsilon \in {}^*\mathbb{R}^+$, y además $\varepsilon \simeq 0$. Luego, $f(0) = 0$ y $f(\varepsilon) = 1$, no estando infinitesimalmente cerca, por lo cual $f(\text{hal}(\varepsilon)) \not\subseteq \text{hal}(f(\varepsilon)) = \text{hal}(1)$. Esto ocurre ya que $\varepsilon \notin \mathbb{R}^+$.

¹¹Notemos que esto es equivalente a probar que la función es continua en 0.

¹²La validez de esta propiedad para números hiperreales se justifica por medio del principio de transferencia.

Para que ocurra esta propiedad para todo $c \in {}^*A$ debe suceder que $f : A \rightarrow \mathbb{R}$ sea uniformemente continua en A . El siguiente teorema exhibe esta caracterización.

Teorema 3.5. *Una función $f : A \rightarrow \mathbb{R}$ es uniformemente continua en A si y solo si para todo par $x, y \in {}^*A$ sucede que $x \simeq y$ implica $f(x) \simeq f(y)$.*

Notemos que, a diferencia del Teorema 3.4, aquí pedimos que la implicación $x \simeq y \implies f(x) \simeq f(y)$ valga para todo $y \in {}^*A$, no solo para $y \in A$.

La demostración del Teorema 3.5 es muy similar a la del Teorema 3.4, con la salvedad de que debemos usar la definición de continuidad uniforme a la hora de trabajar con el principio de transferencia.

Una de las consecuencias más bonitas del Teorema 3.5 es la posibilidad de probar el teorema de Heine-Cantor de una manera muy sencilla.

Teorema 3.6 (Teorema de Heine-Cantor). *Si f es una función real continua en $[a, b] \subseteq \mathbb{R}$, entonces f es uniformemente continua en $[a, b]$.*

Demostración. Comenzaremos tomando $x, y \in {}^*[a, b]$ tales que $x \simeq y$. Como $a \leq x \leq b$, el hiperreal x es limitado. Sea $c := \text{sh}(x)$, como $x \simeq c$ y $a \leq x \leq b$, entonces $a \leq c \leq b$. Así f es continua en c y por lo tanto $f(x) \simeq f(c)$. Como $y \simeq x$ entonces $y \simeq c$ y $f(y) \simeq f(c)$. De esta manera $f(x) \simeq f(y)$. Luego, por teorema 3.5 queda probada la continuidad uniforme de f en $[a, b]$. ■

Observemos que esta demostración es equivalente para cualquier conjunto $B \subseteq \mathbb{R}$ que tenga la propiedad de que “para todo $x \in {}^*B$ existe $c \in B$ tal que $x \simeq c$ ”, o dicho de otra forma “todo punto de *B se encuentra infinitesimalmente cerca de algún punto de B ”. Esto significa que los términos no estándar de *B “no se encuentran muy alejados” de los estándar. De forma que podríamos utilizar solo puntos de B para aproximar, o cubrir, el conjunto *B , es decir

$${}^*B \subseteq \bigcup_{r \in B} \text{hal}(r).$$

En este momento es recomendable que el lector se tome unos minutos para interpretar esta definición y percatarse de la gran y bella coincidencia que existe entre la propiedad que estamos pidiendo que tenga el conjunto B y un tipo particular de conjuntos muy utilizados en análisis. En efecto, esta propiedad resulta ser una caracterización de los conjuntos compactos de $\mathbb{R}$. La demostración nuevamente excede los alcances de esta monografía. Pero consideraré incluir esta reflexión para que no se desperdicie la oportunidad de hablar sobre esta bonita equivalencia.

3.3. Diferenciación

Antes de presentar el teorema que nos brinda la equivalencia entre la definición estándar de diferenciabilidad y su formulación no estándar, resulta conveniente detenernos en una propiedad que pasamos por alto en la sección anterior.

Teorema 3.7. *Si $f : A \rightarrow \mathbb{R}$ es una función real, entonces $\lim_{x \rightarrow c} f(x) = L$ si y solo si $f(x) \simeq L$ para todo $x \in {}^*A$ tal que $x \simeq c$ y $x \neq c$.*

Notemos que la equivalencia es muy similar a la del Teorema 3.4, así como su demostración. Esto no es de sorprender pues los conceptos de límite y continuidad de funciones se relacionan fuertemente.

Ahora si, veamos la equivalencia para la diferenciabilidad:

Teorema 3.8. *Si f es una función real definida en $x \in \mathbb{R}$, entonces $L \in \mathbb{R}$ es la derivada de f en x si y solo si para todo infinitesimal $\varepsilon \neq 0$, se tiene que $f(x + \varepsilon)$ está bien definido y*

$$\frac{f(x + \varepsilon) - f(x)}{\varepsilon} \simeq L.$$

Demostración. Si definimos $g(h) := \frac{f(x+h)-f(x)}{h}$, del Teorema 3.7 se tiene que $\lim_{h \rightarrow 0} g(h) = L$ si y solo si para todo $h \simeq 0$ y $h \neq 0$ sucede que $g(h) \simeq L$. Como $\lim_{h \rightarrow 0} g(h)$ es igual a la derivada de f en x , ya hemos demostrado la equivalencia que deseamos probar. ■

De esta manera, si f es diferenciable en x , entonces

$$f'(x) = \text{sh} \left(\frac{f(x + \varepsilon) - f(x)}{\varepsilon} \right)$$

para todo $\varepsilon \simeq 0$ y $\varepsilon \neq 0$ (es decir, ε infinitesimal y distinto de 0).

Veamos ahora algunos ejemplos de cálculo de derivadas para ciertas funciones.

- Sea $f : \mathbb{R} \rightarrow \mathbb{R}$ definida como $f(x) = x^2$. Luego, para $\varepsilon \simeq 0$ y $\varepsilon \neq 0$

$$\frac{f(x + \varepsilon) - f(x)}{\varepsilon} = \frac{(x + \varepsilon)^2 - x^2}{\varepsilon} = \frac{x^2 + 2\varepsilon x + \varepsilon^2 - x^2}{\varepsilon} = 2x + \varepsilon \simeq 2x.$$

De esta manera, $f'(x) = 2x$ para todo $x \in \mathbb{R}$.

- Consideremos ahora la función $f : \mathbb{R} - \{0\} \rightarrow \mathbb{R}$ dada por $f(x) = \frac{1}{x}$. Luego, si $\varepsilon \simeq 0$ y $\varepsilon \neq 0$ se tiene que

$$\frac{1}{\varepsilon}(f(x + \varepsilon) - f(x)) = \frac{1}{\varepsilon} \left(\frac{1}{x + \varepsilon} - \frac{1}{x} \right) = \frac{1}{\varepsilon} \frac{x - (x + \varepsilon)}{x(x + \varepsilon)} = \frac{1}{\varepsilon} \frac{-\varepsilon}{x^2 + x\varepsilon} = \frac{-1}{x^2 + x\varepsilon}.$$

Como x es limitado, $x\varepsilon \simeq 0$, y así $x^2 + x\varepsilon \simeq x^2$. Además, como $x \neq 0$ ocurre que

$$\frac{-1}{x^2 + x\varepsilon} \simeq \frac{-1}{x^2}.$$

De esta forma, $f'(x) = -\frac{1}{x^2}$ para todo $x \in \mathbb{R} - \{0\}$.

Introduciremos una nueva notación sugestiva. Denotaremos por Δx a una variación o incremento de la variable x , y por Δf al incremento correspondiente a la función f que toma por argumento la variable x . De forma que $\Delta f^{13} := f(x + \Delta x) - f(x)$. Con esta notación podemos reescribir el Teorema 3.8 de la siguiente manera:

L es la derivada de f en x si y solo si $\frac{\Delta f}{\Delta x} \simeq L$ para todo $\Delta x \simeq 0$ y $\Delta x \neq 0$.

¹³Como el valor del incremento de f depende tanto del valor de x como del valor del incremento Δx la manera correcta de denotarlo sería $\Delta f(x, \Delta x)$. Pero para no sobrecargar la notación escribiremos simplemente Δf , entendiendo que su valor puede depender de estos parámetros previamente mencionados.

Por último, definiremos el *diferencial de f en x* asociado a un incremento Δx como $df^{14} := f'(x)\Delta x$. Dado que $f'(x)$ es limitado (en particular es real) se sigue que df es infinitesimal siempre que Δx lo sea. De esta forma, si en lugar de denotar al incremento infinitesimal de x como Δx lo denotáramos por dx , podemos describir la derivada de f en x como

$$f'(x) = \frac{df}{dx}.$$

Tomemos un tiempo para observar como esta expresión no es meramente una “notación”, sino que corresponde al cociente de dos números, ambos infinitesimales. Por comodidad, continuaremos utilizando Δx para referirnos al incremento de x .

Sabemos que si f es diferenciable en x , entonces $\frac{\Delta f}{\Delta x} \simeq f'(x)$, y como $f'(x)$ es un número real se deduce que $\frac{\Delta f}{\Delta x}$ es limitado. Así, $\Delta f = \frac{\Delta f}{\Delta x} \Delta x$, y, si $\Delta x \simeq 0$, se sigue que $\Delta f \simeq 0$. Reescribiendo este último resultado, si $\Delta x \simeq 0$ entonces $f(x + \Delta x) \simeq f(x)$. De esta manera, acabamos de probar que toda función diferenciable en x es continua ese punto.

A simple vista se podría pensar que, para un Δx fijo, Δf y df son lo mismo, pero no es así. Para empezar, df solo está definido cuando f es diferenciable en x . En este caso, $\Delta f \simeq f'(x)\Delta x$, y así $\Delta f \simeq df$. De esta forma, se tiene que $\Delta f - df$ es una cantidad infinitesimal, pero podemos probar algo todavía más fuerte. Como

$$\frac{\Delta f - df}{\Delta x} = \frac{\Delta f}{\Delta x} - \frac{df}{\Delta x} = \frac{\Delta f}{\Delta x} - f'(x) \simeq 0,$$

podemos afirmar que $\Delta f - df$ es *infinitesimal comparado con Δx* .

Si bien no hemos definido previamente el concepto de *infinitesimal comparado*, este resulta bastante intuitivo. Diremos que un número hiperreal a es infinitesimal comparado con otro número hiperreal b si $\frac{a}{b} \simeq 0$. Una manera intuitiva de pensar esto es interpretar que a es de un “orden” o “magnitud” menor que b . Por ejemplo, todo número limitado es infinitesimal comparado con un número ilimitado, y todo infinitesimal es infinitesimal comparado con cualquier número no infinitesimal (valga la redundancia). Además, se tiene que si $\varepsilon \simeq 0$ entonces ε^n es infinitesimal comparado con ε^m siempre que $n > m$.

Todos los resultados mencionados en los párrafos anteriores se pueden resumir en el siguiente teorema:

Teorema 3.9. *Si f es diferenciable en x y $\Delta x \simeq 0$, $x \neq 0$, entonces $\Delta f \simeq 0$, $df \simeq 0$ y existe un infinitesimal $\delta \simeq 0$ (que puede depender de x y Δx) tal que $\Delta f = f'(x)\Delta x + \delta\Delta x$.*

Notemos que el hecho de poder escribir a $\Delta f - df$ como el producto de un infinitesimal δ multiplicado por Δx es equivalente a decir que esta diferencia es infinitesimal comparado con Δx .

Es importante el hecho de que el infinitesimal δ puede depender de Δx . Para que esto no suceda se necesita que f' sea continua en x . No demostraremos este hecho aquí, ya que requeriría un análisis más detallado, pero lo asumiremos como verdadero a partir de este punto.

Otro resultado que sale a la luz de simplemente reordenar los términos es que, si f es diferenciable en x y $\Delta x \simeq 0$, entonces

$$f(x + \Delta x) = f(x) + f'(x)\Delta x + \delta\Delta x.$$

¹⁴De igual manera que ocurre con Δf , la forma más acorde para denotar df es $df(x, \Delta x)$, pues su valor depende tanto de x como de Δx .

Esta resulta ser la expansión en serie de Taylor de primer orden de la función f centrada en x . Notemos que aquí el infinitesimal $\delta\Delta x$ actúa como el resto. De forma que, en lugar de tener una función que al dividirla por Δx tienda a 0, se tiene un infinitesimal que al dividirlo por Δx sigue siendo infinitesimal. Esto es todavía más fuerte cuando f' es continua en x , ya que de esta manera, como se mencionó anteriormente, el infinitesimal δ solo depende de x y no de Δx . Por lo que el resto será linealmente proporcional al incremento Δx elegido. Si denotamos por $\ell(\Delta x) = f(x) + f'(x)\Delta x$, entonces $f(x + \Delta x) \simeq \ell(\Delta x)$, y se tiene la aproximación lineal de f alrededor de x .

Veamos ahora como obtener las reglas básicas de derivación a partir de argumentos no estándar.

Proposición 3.10. *Si f y g son diferenciables en x , entonces $(f + g)'(x) = f'(x) + g'(x)$.*

Demostración. Sea $\Delta x \simeq 0$ y $\Delta x \neq 0$. Como f y g son diferenciables en x , entonces $f(x + \Delta x)$ y $g(x + \Delta x)$ se encuentran definidos. Luego

$$\Delta(f + g) = (f + g)(x + \Delta x) - (f + g)(x) = f(x + \Delta x) + g(x + \Delta x) - f(x) - g(x).$$

Dividiendo por Δx obtenemos que

$$\frac{\Delta(f + g)}{\Delta x} = \frac{f(x + \Delta x) - f(x)}{\Delta x} + \frac{g(x + \Delta x) - g(x)}{\Delta x} = \frac{\Delta f}{\Delta x} + \frac{\Delta g}{\Delta x} \simeq f'(x) + g'(x).$$

Finalmente, por Teorema 3.8 se tiene que $(f + g)'(x) = f'(x) + g'(x)$. ■

Proposición 3.11. *Si f y g son diferenciables en x , entonces $(fg)'(x) = f'(x)g(x) + f(x)g'(x)$.*

Demostración. Sea $\Delta x \simeq 0$ y $\Delta x \neq 0$. Luego,

$$\Delta(fg) = (fg)(x + \Delta x) - (fg)(x) = f(x + \Delta x)g(x + \Delta x) - f(x)g(x).$$

Como $f(x + \Delta x) = f(x) + \Delta f$ y $g(x + \Delta x) = g(x) + \Delta g$ se sigue que

$$\Delta(fg) = (f(x) + \Delta f)(g(x) + \Delta g) - f(x)g(x) = \Delta f \cdot g(x) + f(x) \cdot \Delta g + \Delta f \cdot \Delta g.$$

Dividiendo por Δx obtenemos

$$\frac{\Delta(fg)}{\Delta x} = \frac{\Delta f}{\Delta x}g(x) + f(x)\frac{\Delta g}{\Delta x} + \frac{\Delta f}{\Delta x}\Delta g.$$

Como f es diferenciable en x , $\frac{\Delta f}{\Delta x}$ es limitado, y como g es diferenciable en x , $\Delta g \simeq 0$. De esta manera se tiene que $\frac{\Delta f}{\Delta x}\Delta g \simeq 0$. Así

$$\frac{\Delta(fg)}{\Delta x} \simeq f'(x)g(x) + f(x)g'(x).$$

Finalmente, por el Teorema 3.8 concluimos que $(fg)'(x) = f'(x)g(x) + f(x)g'(x)$. ■

Proposición 3.12. *Si f es diferenciable en x y g es diferenciable en $f(x)$, entonces $g \circ f$ es diferenciable en x . Mas aún, $(g \circ f)'(x) = g'(f(x))f'(x)$.*

Demostración. Sea $\Delta x \simeq 0$ y $\Delta x \neq 0$. Como f es diferenciable en x , entonces $f(x + \Delta x)$ está bien definido y $f(x + \Delta x) \simeq f(x)$. De la misma manera, como g es diferenciable en $f(x)$ y $\Delta f \simeq 0$, entonces $g(f(x) + \Delta f)$ está bien definido y

$$(g \circ f)(x + \Delta x) = g(f(x + \Delta x)) = g(f(x) + \Delta f) \quad (1)$$

existe. Ahora, notemos que

$$\Delta(g \circ f)(x, \Delta x) = (g \circ f)(x + \Delta x) - (g \circ f)(x) = g(f(x + \Delta x)) - g(f(x)).$$

Pero por (1) se sigue

$$\Delta(g \circ f)(x, \Delta x) = g(f(x) + \Delta f) - g(f(x)) = \Delta g(f(x), \Delta f).$$

Recordemos que por el Teorema 3.9 existe $\delta \simeq 0$ tal que

$$\Delta g(f(x), \Delta f) = g'(f(x))\Delta f + \delta \Delta f.$$

Así, si $\Delta f = 0$ entonces $\Delta(g \circ f)(x, \Delta x) = \Delta g(f(x), \Delta f) = 0$. Además, $f'(x) = 0$ y en consecuencia $(g \circ f)'(x) = g'(f(x))f'(x) = 0$, valiendo el resultado del teorema. Supongamos entonces que $\Delta f \neq 0$. Dividiendo por Δx en la ecuación (1) se sigue que

$$\frac{\Delta(g \circ f)(x, \Delta x)}{\Delta x} = \frac{\Delta g(f(x), \Delta f)}{\Delta x} = \frac{\Delta g(f(x), \Delta f)}{\Delta f} \cdot \frac{\Delta f}{\Delta x}.$$

Luego, como

$$\frac{\Delta g(f(x), \Delta f)}{\Delta f} \simeq g'(f(x)) \quad \text{y} \quad \frac{\Delta f}{\Delta x} \simeq f'(x),$$

se tiene que

$$\frac{\Delta(g \circ f)}{\Delta x} \simeq g'(f(x))f'(x).$$

Finalmente, por el Teorema 3.8 concluimos que $(g \circ f)'(x) = g'(f(x))f'(x)$. ■

Antes de pasar a la siguiente sección mencionaremos un resultado breve pero interesante, que muestra cómo se pueden interpretar las derivadas como una suerte de “sombras” de los operadores de diferencia finita.

Para una función f y $h \in \mathbb{R} - \{0\}$ definimos el *operador diferencia finita* como

$$D_h[f](x) := \frac{f(x + h) - f(x)}{h}.$$

Llamaremos h al *paso* de la diferencia finita. Dado $n \in \mathbb{N}$, definimos el *operador diferencia finita de orden n* para $n = 1$ como $D_h^1 = D_h$, y para $n > 1$ de manera recursiva como

$$D_h^n[f](x) := D_h[D_h^{n-1}[f]](x).$$

Es posible obtener una fórmula cerrada para este operador de orden n , pero no nos interesa para lo que vamos a realizar. Otras observaciones pertinentes es que este operador es lineal

$$D_h^n[f + \lambda g](x) = D_h^n[f](x) + \lambda D_h^n[g](x)$$

para todo $\lambda \in \mathbb{R}$, $n \in \mathbb{N}$ y cualesquiera funciones f y g . Además, si evaluamos el operador en una función constante $p(x) = k$ se tiene que $D_h^n[p](x) = 0$.

Si bien hemos definido este operador para $h \in \mathbb{R} - \{0\}$ puede ser usado sin ningún problema para $h \in {}^*\mathbb{R} - \{0\}$ (siempre que la función lo permita). Notemos que, si tomamos $\varepsilon \simeq 0$, $\varepsilon \neq 0$ y f diferenciable en x entonces $\text{sh}(D_\varepsilon[f](x)) = f'(x)$, es decir,

$$D_\varepsilon[f](x) \simeq f'(x).$$

Esto mismo se puede extender a operadores de diferencia finita de orden n

Teorema 3.13. *Si f es n veces diferenciable en x , entonces para todo $\varepsilon \simeq 0$ y $\varepsilon \neq 0$ sucede que*

$$D_\varepsilon^n[f](x) \simeq f^{(n)}(x).$$

Demostración. Realizaremos la demostración por inducción. El paso base $n = 1$ ya ha sido probado en el párrafo previo al enunciado del teorema.

Supongamos ahora que para $n \in \mathbb{N}$ vale $D_\varepsilon^n[f](x) \simeq f^{(n)}(x)$, veamos ahora que esto ocurre también para $n + 1$.

Dado un ε fijo se tiene que, en virtud de la hipótesis inductiva, $D_\varepsilon^n[f](x) - f^{(n)}(x) = \delta$ un infinitesimal constante. De esta manera, por linealidad del operador de diferencia finita,

$$D_\varepsilon^{n+1}[f](x) - D_\varepsilon[f^{(n)}](x) = D_\varepsilon[D_\varepsilon^n[f] - f^{(n)}](x) = D_\varepsilon[\delta](x) = 0.$$

Así, obtenemos que $D_\varepsilon^n[f](x) = D_\varepsilon[f^{(n)}](x)$. Luego, por el paso base se sigue que

$$D_\varepsilon[f^{(n)}](x) \simeq (f^{(n)})'(x),$$

y como $(f^{(n)})'(x) = f^{(n+1)}(x)$ por definición, concluimos finalmente que

$$D_\varepsilon^{n+1}[f](x) \simeq f^{(n+1)}(x),$$

para todo $\varepsilon \simeq 0$ y $\varepsilon \neq 0$. ■

Este resultado establece que, al calcular la diferencia finita de orden n de una función f en el punto x con un paso infinitesimal ε , se obtiene un valor infinitesimalmente cercano a la n -ésima derivada de f en x . En otras palabras, podemos aproximar el valor de $f^{(n)}(x)$ mediante $D_\varepsilon^n[f](x)$, siempre que el paso ε sea suficientemente pequeño.

Vale la pena aclarar que el error $\xi = D_\varepsilon^n[f](x) - f^{(n)}(x)$, si bien es infinitesimal, puede depender del valor de x y ε . Si además suponemos que $f^{(n)}$ es continua en x , entonces el valor de ξ no dependerá de ε . Este hecho está estrechamente relacionado con que el error de aproximación, al estimar derivadas mediante operadores de diferencias finitas de este tipo, sea de orden lineal.

3.4. Integración

En esta sección desarrollaremos los conceptos básicos relacionados a integración con la mirada del análisis no estándar.

Comenzaremos definiendo que entendemos por *partición*. Dado intervalo cerrado $[a, b]$, una partición de este intervalo es un conjunto finito de puntos $P\{x_0, \dots, x_n\}$ que cumple $a = x_0 < x_1 < \dots < x_{n-1} < x_n = b$. Para nuestros propósitos, será más cómodo trabajar

con un tipo particular de particiones. Dado un número real positivo Δx , podemos construir una *partición uniforme* de la forma $P_{\Delta x} = \{x_0, \dots, x_n\}$ donde $x_k = a + k\Delta x = x_{k-1} + \Delta x$ para $k = 1, \dots, n$, y n es el menor número natural tal que $a + n\Delta x \geq b$. Además, definimos la longitud del i -ésimo intervalo como $\Delta x_i = x_i - x_{i-1}$ para $i = 1, \dots, n$. Si nuestra partición es de la forma $P_{\Delta x}$ es evidente que $\Delta x_i = \Delta x$ para todo i .

Sea $f : [a, b] \rightarrow \mathbb{R}$ una función acotada y $P = \{x_0, \dots, x_n\}$ una partición del intervalo $[a, b]$. Para cada subintervalo $[x_{i-1}, x_i]$ definimos

$$M_i := \sup_{x \in [x_{i-1}, x_i]} \{f(x)\}, \quad \text{y} \quad m_i := \inf_{x \in [x_{i-1}, x_i]} \{f(x)\}.$$

En base a esto definimos:

- la *suma superior de Riemann* $U_a^b(f, P) = \sum_{i=1}^n M_i \Delta x_i$,
- la *suma inferior de Riemann* $L_a^b(f, P) = \sum_{i=1}^n m_i \Delta x_i$,
- la *suma ordinaria de Riemann* $S_a^b(f, P) = \sum_{i=1}^n f(x_{i-1}) \Delta x_i$.

Si tenemos una partición de la forma $P_{\Delta x}$, está claro que la dependencia de las distintas sumas con respecto a la partición se traduce directamente en una dependencia de Δx . De ahora en adelante trabajaremos con este tipo de particiones, y simplemente escribiremos $U_a^b(f, \Delta x)$, $L_a^b(f, \Delta x)$ y $S_a^b(f, \Delta x)$. Por lo tanto, para una función f dada y un intervalo $[a, b]$ fijo, todas las sumas anteriores resultan ser simplemente funciones reales que dependen de la variable Δx . De esta manera, las sumas pueden entenderse, para valores de Δx infinitesimales positivos, como la extensión directa de una función.

Uno de los primeros resultados que se pueden probar es que, para todo Δx infinitesimal positivo, se cumple $L_a^b(f, \Delta x) \simeq U_a^b(f, \Delta x)$. Además, para todo par de infinitesimales positivos Δx_1 y Δx_2 se tiene

$$U_a^b(f, \Delta x_1) \simeq U_a^b(f, \Delta x_2), \quad L_a^b(f, \Delta x_1) \simeq L_a^b(f, \Delta x_2), \quad S_a^b(f, \Delta x_1) \simeq S_a^b(f, \Delta x_2).$$

Dado que todas las sumas son acotadas (por el principio de transferencia), se concluye que las mismas son números limitados. De esta manera, podemos definir

$$I_a^b(f) := \text{sh}(S_a^b(f, \Delta x))$$

para todo $\Delta x \simeq 0$ con $\Delta x > 0$. Finalmente, como la definición estándar de la integral de f entre a y b resulta ser el límite de las sumas de Riemann (cuando este existe), resulta evidente que

$$\int_a^b f(x) dx = I_a^b(f) = \text{sh}(S_a^b(f, \Delta x)),$$

para todo $\Delta x \simeq 0$ con $\Delta x > 0$.

Esta es la caracterización no estándar de las integrales. Es cierto que no resulta tan impactante como las que equivalencias que vimos en las secciones previas. Para tener una intuición más satisfactoria de la integral desde el análisis no estándar se requiere profundizar mucho más en el tema. Incluso se deben estudiar otras formulaciones y construcciones del análisis no estándar —como la construcción de Nelson¹⁵ a través de conjuntos internos,

¹⁵Al comienzo de la Sección 4.1 del Capítulo 4 se desarrolla brevemente la diferencia entre esta construcción y la utilizada en la monografía.

que se puede ver en [7]—, que permiten llegar a expresiones del tipo: “la integral es una suma ilimitada de elementos infinitesimales”.

A pesar de esto, existen integrales definidas que podemos calcular de manera muy sencilla por medio de esta caracterización. A continuación mostraremos algunos ejemplos.

- Consideremos la función identidad $f(x) = x$, y calculemos el valor de la integral definida en el intervalo $[0, 1]$. Para ello tomaremos una partición uniforme

$$P_n = \{x_0, \dots, x_n\}, \quad \text{con } x_i = \frac{i}{n} \text{ para } i = 0, \dots, n,$$

con $n \in \mathbb{N}$. De esta forma, la suma ordinaria de Riemann resulta

$$S_0^1\left(f, \frac{1}{n}\right) = \sum_{i=1}^n \frac{i}{n} \frac{1}{n} = \frac{1}{n^2} \sum_{i=1}^n i = \frac{1}{n^2} \frac{n(n+1)}{2} = \frac{n^2 + n}{2n^2} = \frac{1}{2} + \frac{1}{2n}.$$

Si consideramos ahora $N \in {}^*\mathbb{N}_\infty$, en virtud del principio de transferencia

$$S_0^1\left(f, \frac{1}{N}\right) = \frac{1}{2} + \frac{1}{2N} \simeq \frac{1}{2},$$

pues $\frac{1}{2N}$ es infinitesimal. De esta manera, hemos probado que

$$\int_0^1 x \, dx = \text{sh}\left(S_0^1\left(x, \frac{1}{N}\right)\right) = \frac{1}{2}.$$

- Tomemos ahora la función $f(x) = x^2$ y el intervalo $[0, 1]$. Utilizaremos la misma partición uniforme del caso anterior, de forma que

$$S_0^1\left(f, \frac{1}{n}\right) = \sum_{i=1}^n \left(\frac{i}{n}\right)^2 \frac{1}{n} = \frac{1}{n^3} \sum_{i=1}^n i^2 = \frac{1}{n^3} \frac{n(n+1)(2n+1)}{6} = \frac{1}{3} + \frac{1}{2n} + \frac{1}{6n^2}.$$

Luego, tomamos $N \in {}^*\mathbb{N}_\infty$, y del mismo modo que el ejemplo anterior se sigue que

$$\int_0^1 x^2 \, dx = \text{sh}\left(S_0^1\left(f, \frac{1}{N}\right)\right) = \frac{1}{3}.$$

- Como último ejemplo veamos el valor de la integral de la función $f(x) = \frac{1}{x}$ en el intervalo $[1, e]$. Tomaremos una partición del intervalo de la forma

$$P_n^{16} = \{x_0, \dots, x_n\}, \quad \text{con } x_i = e^{i/n} \text{ para } i = 0, \dots, n,$$

con $n \in \mathbb{N}$. Así, $\Delta x_i = e^{i/n} - e^{(i-1)/n}$. De esta manera,

$$S_1^e(f, P_n) = \sum_{i=1}^n \frac{1}{e^{(i-1)/n}} (e^{i/n} - e^{(i-1)/n}) = \sum_{i=1}^n (e^{1/n} - 1) = n (e^{1/n} - 1).$$

Sea $N \in {}^*\mathbb{N}_\infty$. Recordemos que, al ser N ilimitado, se tiene $(1 + \frac{1}{N})^N \simeq e$. Así,

$$S_1^e(f, P_N) = N (e^{1/N} - 1) \simeq N \left(\left(1 + \frac{1}{N}\right)^{N/N} - 1 \right) = N \frac{1}{N} = 1.$$

Concluyendo finalmente que

$$\int_1^e \frac{1}{x} \, dx = 1.$$

¹⁶Si bien se utilizó la misma notación, esta partición no resulta una partición uniforme.

Veamos ahora cómo demostrar los teoremas fundamentales del cálculo.

Teorema 3.14. *Sea $f : [a, b] \rightarrow \mathbb{R}$ una función real continua. Entonces la función $F(x) = \int_a^x f(t) dt$ es diferenciable en $[a, b]$, y cumple con $F'(x) = f(x)$ para todo $x \in [a, b]$.*

Demostración. Sea $x \in [a, b]$ fijo. Si tomamos $\Delta x \in \mathbb{R}^+$ tal que $\Delta x < b - x$ se sigue

$$F(x + \Delta x) - F(x) = \int_x^{x+\Delta x} f(t) dt.$$

Como f es continua en $[x, x + \Delta x]$, f alcanza un máximo y un mínimo en x_1 y x_2 respectivamente. Luego,

$$\begin{aligned} f(x_2)\Delta x &\leq \int_x^{x+\Delta x} f(t) dt \leq f(x_1)\Delta x, \\ f(x_2) &\leq \frac{1}{\Delta x} \int_x^{x+\Delta x} f(t) dt \leq f(x_1). \end{aligned} \quad (2)$$

Por el principio de transferencia, si $\Delta x \in {}^*\mathbb{R}^+$ (en particular $\Delta x \simeq 0$ y $\Delta x > 0$), deben existir hiperreales $x_1, x_2 \in [x, x + \Delta x]$ que cumplan con la misma desigualdad (2). Como, $x \simeq \Delta x$ entonces $x_1 \simeq x \simeq x_2$, y por continuidad de f , $f(x_1) \simeq f(x) \simeq f(x_2)$. De esta manera, para todo infinitesimal positivo Δx se tiene

$$\frac{F(x + \Delta x) - F(x)}{\Delta x} = \frac{1}{\Delta x} \int_x^{x+\Delta x} f(t) dt \simeq f(x).$$

De manera análoga, puede demostrarse el mismo resultado para Δx infinitesimal negativo, completando así la prueba de que $F'(x) = f(x)$. ■

Teorema 3.15. *Sea $G : [a, b] \rightarrow \mathbb{R}$ una función continua tal que $G'(x) = f(x)$ para todo $x \in [a, b]$. Entonces $\int_a^b f(x) = G(b) - G(a)$.*

Para demostrar este segundo teorema no es necesario requerir a argumentos no estándar, basta con utilizar el Teorema 3.14.

Para concluir la sección de integrales, mostraremos cómo deducir la fórmula del promedio integral de una función a partir del promedio de cantidades finitas.

Sea f es una función continua en $[a, b]$, definimos

$$\text{Av}(n) := \frac{f(x_0) + \cdots + f(x_n)}{n},$$

con $x_i = a + i \frac{b-a}{n}$ para $i = 0, \dots, n$. A continuación, veremos que

$$\text{Av}(N) \simeq \frac{1}{b-a} \int_a^b f(x) dx$$

para todo $N \in {}^*\mathbb{N}_\infty$. Lo que indicaría que la extensión de la “función promedio” evaluada en hipernaturales ilimitados se encuentra infinitesimalmente cerca del promedio integral.

Primero definimos $\Delta x = \frac{b-a}{n}$, así $x_i = a + i\Delta x$. Luego,

$$Av(n) = \sum_{i=1}^n f(x_{i-1}) \frac{\Delta x}{b-a} = \frac{1}{b-a} \sum_{i=1}^n f(x_{i-1}) \Delta x = \frac{1}{b-a} S_a^b(f, \Delta x).$$

Notemos que, si $N \in {}^*\mathbb{N}_\infty$, entonces $\Delta x = \frac{b-a}{N} \simeq 0$. Por lo tanto

$$Av(N) = \frac{1}{b-a} S_a^b(f, \Delta x) \simeq \frac{1}{b-a} \int_a^b f(x) dx.$$

Concluyendo así la demostración.

4. Aplicaciones

“Eso no lo entendí hasta que una visión de la noche me reveló que también son catorce [son infinitos] los mares y los templos.”

— Jorge Luis Borges, *La casa de Asterion*

En este último capítulo se busca abordar algunas aplicaciones de todo lo que hemos trabajado a lo largo de esta monografía. En primer lugar, presentaremos una interpretación del teorema de existencia para ecuaciones diferenciales ordinarias desde la perspectiva del análisis no estándar. Posteriormente, mostraremos cómo, mediante el uso de cantidades infinitesimales, puede deducirse que el movimiento transversal de una cuerda vibrante satisface la ecuación de onda. Por último, comentaremos algunas aplicaciones interesantes que no entraron en la monografía.

4.1. Ecuaciones diferenciales ordinarias

A continuación, presentaremos algunos resultados del capítulo titulado *Applications of Nonstandard Analysis in Ordinary Differential Equations* de [1]. En dicho capítulo se demuestran diversas proposiciones relacionadas con ecuaciones diferenciales ordinarias mediante el uso del análisis no estándar. En particular, enunciaremos una forma más débil del Teorema 3.1 (de dicho capítulo), lo cual nos permitirá mostrar aplicaciones prácticas sin requerir todo el trasfondo teórico necesario para la formulación general del resultado. Debido a que el capítulo está redactado utilizando nomenclatura y definiciones propias del “lenguaje de Nelson”¹⁷, me he visto en la necesidad de reescribir los teoremas y resultados para adecuarlos al lenguaje utilizado en esta monografía.

Un primer concepto que resulta importante aclarar es el de *S-continuidad*. Diremos que una función $f : {}^*A \rightarrow {}^*\mathbb{R}$ es *S-continua* si para todo $c \in A$ y $x \in {}^*A$ se tiene que si $x \simeq c$ entonces $f(x) \simeq f(c)$. Notemos la similitud entre esta definición y el Teorema 3.4. En particular, si f es la extensión de una función real con dominio A , este teorema afirma que dicha función real es continua en el sentido estándar. Sin embargo, la propiedad de *S-continuidad* es mucho más general, puede aplicarse incluso a funciones hiperreales que no sean la extensión de ninguna función real (objetos con los que todavía no hemos trabajado).

Enunciaremos ahora un teorema que nos permitirá obtener información acerca de existencia y unicidad de ecuaciones diferenciales ordinarias de primer orden con condiciones iniciales.

¹⁷Edward Nelson fue un matemático que desarrolló la teoría de conjuntos internos (IST por sus siglas en inglés), una fundamentación axiomática para una parte del análisis no estándar de Robinson. La diferencia más notable entre ambos enfoques es que, de forma coloquial, mientras que Robinson construye nuevas entidades “extendiendo” las usuales, Nelson considera al conjunto de los números hiperreales como el conjunto natural con el que siempre se trabajó, agregando axiomas que permiten distinguir las entidades “estándar” dentro de este conjunto, las cuales corresponden a los números reales.

Teorema 4.1. Sea $f : \mathbb{R} \rightarrow \mathbb{R}$ una función real y acotada, y sea $\{t_i\}_{i=0,1,\dots,N}$ una hipersucesión de números hiperreales tales que

$$t_0 < t_1 < \dots < t_N, \quad t_i \simeq t_{i+1} \quad \text{para } i = 0, 1, \dots, N-1, \quad t_0 \in \mathbb{L}, \quad t_N - t_0 \notin \mathbb{L}.$$

Además, sea $\{\xi_i\}_{i=0,1,\dots,N}$ una hipersucesión de números hiperreales tal que

$$\xi_0 \in \mathbb{L}, \quad \frac{\xi_{i+1} - \xi_i}{t_{i+1} - t_i} \simeq f(\xi_i) \quad \text{para } i = 0, 1, \dots, N-1.$$

Entonces, la función $h : [t_0, t_N) \rightarrow {}^*\mathbb{R}$ definida como $h(t) := \xi_i$ para $t \in [t_i, t_{i+1})$ es S -continua, y la función x definida como $x(t) := \text{sh}(h(t))$ para todo $t \in [\text{sh}(t_0), \text{sh}(t_N))$ es solución del problema a valores iniciales (PVI)

$$\begin{cases} \dot{x} = f(x), \\ x(\text{sh}(t_0)) = \text{sh}(\xi_0), \end{cases}$$

donde $\dot{x}$ representa $\frac{dx}{dt}$, la derivada de x respecto a t .

Notemos que, en ambas hipersucesiones, el número N puede ser un número hipernatural ilimitado a priori. De hecho, debe serlo, ya que de lo contrario no podría ocurrir que $t_i \simeq t_{i+1}$ y $t_N \not\simeq t_0$. Recordemos que, a pesar del nombre, resulta conveniente pensar las hipersucesiones como funciones de ${}^*\mathbb{N}$ en ${}^*\mathbb{R}$, y no como sucesiones bajo la perspectiva estándar.

La demostración completa de este teorema excede los alcances de la monografía, pero mencionaremos como se realiza. Lo primero es demostrar que la función h es S -continua, para ello se usan las propiedades de los elementos de las hipersucesiones y el hecho de que f es acotada. Luego, como $h(\text{sh}(t_0)) = \xi_0$ es limitado y el intervalo $[t_0, t_N)$ es lo que se denomina un *conjunto interno*, se utiliza un resultado denominado *teorema de la sombra continua* (*continuous shadow theorem*) para afirmar la existencia y continuidad de una función x continua que satisface los requisitos del teorema. Por último, se comprueba que la integral de $f(x)$ es la función x , por medio del teorema fundamental del cálculo.

Como corolario del teorema anterior podemos afirmar la existencia de soluciones para un problema de valores iniciales.

Corolario 4.1.1. Sea $f : \mathbb{R} \rightarrow \mathbb{R}$ una función continua, entonces el problema de valores iniciales

$$\begin{cases} \dot{x} = f(x), \\ x(t_0) = x_0, \end{cases}$$

tiene al menos una solución $x : I \rightarrow \mathbb{R}$ tal que $x(t_0) = x_0$ y $\dot{x}(t) = f(x(t))$ para todo $t \in I$. Donde I es un intervalo máximo.

Demostración. Sea $\varepsilon \simeq 0$ y $\varepsilon > 0$, y sea $N \in {}^*\mathbb{N}_\infty$ dado. Para cada $i = 0, 1, \dots, N$ definimos

$$t_i = t_0 + \varepsilon i, \quad \xi_0 = x_0, \quad \xi_{i+1} = \xi_i + \varepsilon f(\xi_i).$$

De esta forma, tanto $\{t_i\}$ como $\{\xi_i\}$ cumplen con las hipótesis del Teorema 4.1, obteniendo así que la función x definida sobre el intervalo $[t_0, \text{sh}(t_N))$ es una solución del problema de valores iniciales. Si quisiéramos probar la existencia de una solución para un intervalo de la forma $(a, t_0]$ debemos tomar ε infinitesimal negativo y elegir un infinitesimal ilimitado $M \in {}^*\mathbb{N}_\infty$ tal que $\text{sh}(t_M) = \text{sh}(t_0 + \varepsilon M) = a$. ■

Un lector que haya realizado un curso de análisis numérico podrá reconocer que esta forma de construir la sucesión $\{\xi_i\}$ es similar al método de Euler. La diferencia aquí radica en que los pasos del intervalo de la variable t resultan ser infinitesimales, y en consecuencia, la cantidad de elementos de la hipersucesión $\{\xi_i\}$ que aproximará la función h es una cantidad hipernatural ilimitada. Una manera intuitiva de interpretar esto es considerar que, si se refina lo suficiente el intervalo de la variable t , de modo que los incrementos sean infinitesimales, los valores obtenidos mediante el método de Euler estarán infinitesimalmente cerca a los valores de la función solución. Bajo este criterio, es posible emplear cualquier método numérico para aproximar soluciones de PVI a la hora de definir los ξ_i .

Para finalizar esta sección, veamos un ejemplo de cómo obtener una solución a un PVI por medio de este método. Pensemos en el siguiente problema de valores iniciales

$$\begin{cases} \dot{x}(t) &= x(t), \\ x(0) &= 1. \end{cases}$$

Aquí $t_0 = 0$, $x_0 = 1$ y f es la función identidad. Con esto en mente, y acorde al Corolario 4.1.1, definimos

$$t_i = \varepsilon i, \quad \xi_0 = 1, \quad \xi_{i+1} = \xi_i + \varepsilon \xi_i = (1 + \varepsilon) \xi_i.$$

Es fácil ver que, por inducción, podemos escribir

$$\xi_i = (1 + \varepsilon)^i \xi_0 = (1 + \varepsilon)^i$$

para todo $i \in {}^*\mathbb{N}$, debido al principio de transferencia. De esta manera, la sombra de la función h , definida como $h(t) = \xi_i = (1 + \varepsilon)^i$ para $t \in [\varepsilon i, \varepsilon(i + 1))$, es una solución del PVI. Es posible mostrar, por medio de argumentos no estándar, que

$$(1 + \varepsilon)^{1/\varepsilon} \simeq e$$

para todo $\varepsilon \simeq 0$ y $\varepsilon \neq 0$. Así, para $t \in [\varepsilon i, \varepsilon(i + 1))$, se tiene

$$h(t) = \xi_i = (1 + \varepsilon)^i = ((1 + \varepsilon)^{1/\varepsilon})^{\varepsilon i} \simeq e^{\varepsilon i} \simeq e^t.$$

Como e^t es real, se sigue que $\text{sh}(h(t)) = e^t$. De esta manera, e^t con $t \in [0, \text{sh}(t_N))$ es una solución al PVI, para algún $N \in {}^*\mathbb{N}_\infty$ elegido. Notemos que, dado que la elección de N es arbitraria, se puede concluir que esta solución vale para $t \in [0, \infty)$. Análogamente, podemos probar la misma también es válida para $t \in (-\infty, 0]$. De este modo e^t resulta ser una solución del PVI para todo $t \in \mathbb{R}$.

4.2. La ecuación de onda en una cuerda

Presentaremos una deducción clásica en cualquier curso básico de física¹⁸: el movimiento transversal de una cuerda vibrante satisface la ecuación de onda, pero utilizando análisis no estándar. Denotaremos por x a la posición horizontal de la cuerda, por y a la posición vertical, por T a la magnitud de la tensión en la cuerda y por θ al ángulo que forma la tensión con la horizontal. Las variables y , T y θ serán funciones de la posición horizontal x y del tiempo t . Para simplificar la notación, realizaremos el análisis a un tiempo fijo, de modo que estas variables dependerán únicamente de x .

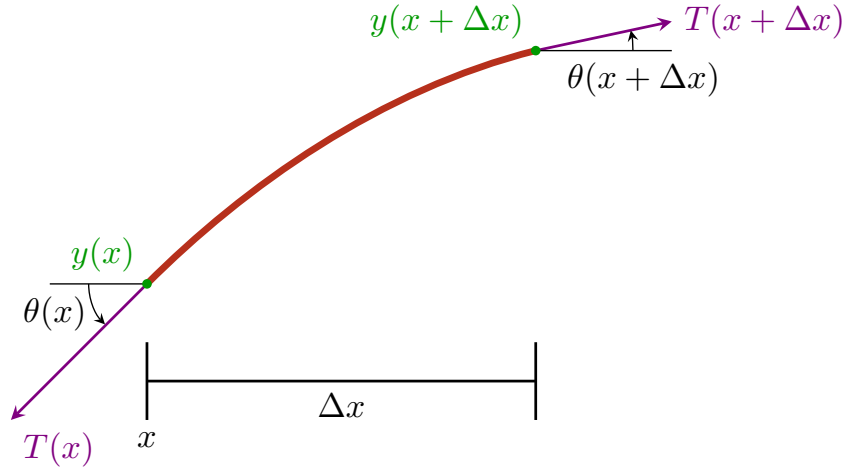


Figura 1: Esquema de la cuerda con tensiones.

Consideremos un pequeño segmento de cuerda con una separación horizontal $\Delta x \neq 0$, como se muestra en la Figura 1. Notemos que Δx , en principio, no tiene por qué ser infinitesimal.

Si interpretamos a la derivada de y en x como la pendiente de la recta tangente a la función y en el punto $(x, y(x))$, se tiene que $\tan(\theta(x)) = y'(x)$, y en consecuencia

$$\sin(\theta(x)) = y'(x) \cos(\theta(x)). \quad (3)$$

De manera análoga, se sigue que

$$\sin(\theta(x + \Delta x)) = y'(x + \Delta x) \cos(\theta(x + \Delta x)). \quad (4)$$

Debido a que la cuerda se encuentra en equilibrio de fuerzas horizontales, aplicando la segunda ley de Newton obtenemos

$$T(x + \Delta x) \cos(\theta(x + \Delta x)) = T(x) \cos(\theta(x)). \quad (5)$$

Luego, si denotamos por m la masa y por a la aceleración vertical del segmento de cuerda, aplicando nuevamente la segunda ley de Newton se tiene

$$ma = T(x + \Delta x) \sin(\theta(x + \Delta x)) - T(x) \sin(\theta(x)).$$

Notemos que la aceleración a corresponde a la aceleración vertical total de la cuerda, evaluada en su centro de masa. En otras palabras, $a = a(\tilde{x})$ para algún $x \leq \tilde{x} \leq x + \Delta x$. Por consecuencia de (3) y (4) se tiene que

$$ma = T(x + \Delta x) y'(x + \Delta x) \cos(\theta(x + \Delta x)) - T(x) y'(x) \cos(\theta(x)),$$

y utilizando la igualdad (5) obtenemos

$$ma = (y'(x + \Delta x) - y'(x)) T(x) \cos(\theta(x)). \quad (6)$$

¹⁸Véase la Sección 16.6 del libro de Serway y Jewett [9] para una deducción estándar.

Podemos calcular la masa m de segmento de cuerda por medio de la integral

$$m = \int_x^{x+\Delta x} \rho(s) \sqrt{1 + (y'(s))^2} ds = \int_x^{x+\Delta x} \rho(s) \cos(\theta(s)) ds. \quad (7)$$

De esta forma, reemplazando (7) en (6) y dividiendo ambos lados por Δx se tiene

$$\frac{1}{\Delta x} \int_x^{x+\Delta x} \rho(s) \cos(\theta(s)) ds \cdot a = \frac{y'(x + \Delta x) - y'(x)}{\Delta x} T(x) \cos(\theta(x)). \quad (8)$$

Ahora, si tomamos $\Delta x \simeq 0$, entonces

$$\text{sh} \left(\frac{y'(x + \Delta x) - y'(x)}{\Delta x} \right) = y''(x),$$

y en virtud del teorema fundamental del cálculo

$$\text{sh} \left(\frac{1}{\Delta x} \int_x^{x+\Delta x} \rho(s) \cos(\theta(s)) ds \right) = \rho(x) \cos(\theta(x)).$$

Además, como $x \leq \tilde{x} \leq x + \Delta x$, se tiene que $x \simeq \tilde{x}$. Bajo la suposición de continuidad de a se sigue que $a(\tilde{x}) \simeq a(x)$ y, al ser este último real, $\text{sh}(a(\tilde{x})) = a(x)$. Finalmente, utilizando estas igualdades en la ecuación (8) obtenemos que

$$\rho(x) \cos(\theta(x)) a(x) = y''(x) \cos(\theta(x)).$$

Como $y'(x)$ es continua, se tiene que $\cos(\theta(x)) \neq 0$ ¹⁹, y en consecuencia $\rho(x) a(x) = y''(x) T(x)$. Finalmente, reordenando los términos obtenemos

$$a(x) = \frac{T(x)}{\rho(x)} y''(x),$$

y escribiéndolo en un lenguaje más familiar

$$\frac{\partial^2 y}{\partial t^2}(x) = \frac{T(x)}{\rho(x)} \frac{\partial^2 y}{\partial x^2}(x).$$

De esta manera, hemos obtenido la ecuación de onda para una cuerda vibrante.

A simple vista la deducción no difiere mucho de la que se puede encontrar en textos de física. La principal diferencia radica en que, en lugar de realizar un paso al límite a la hora de dividir por Δx , obtenemos los resultados deseados a partir de considerar a Δx infinitesimal y aplicar argumentos no estándar.

4.3. Otras aplicaciones

En esta sección comentaremos brevemente otro posible uso del análisis no estándar en la física, en particular en la mecánica analítica.

En la Sección 4.8 del libro de Goldstein [5] se introduce el concepto de *rotaciones infinitesimales*, como aproximación lineal a las rotaciones finitas, un concepto clave a la hora

¹⁹Recordemos que $y'(x) = \tan(\theta(x))$.

de estudiar la cinemática y mecánica de cuerpos rígidos. En este contexto, las rotaciones infinitesimales son generadas por matrices de rotación, cuyos elementos, en lugar de ser senos y cosenos, corresponden a las aproximaciones de estas funciones para ángulos pequeños. A lo largo de dicha sección se utilizan argumentos usuales del cálculo estándar para justificar estas aproximaciones, como la omisión de términos infinitesimales de orden superior debido a ser “muy pequeños en comparación al resto”. Esta práctica sugiere la posibilidad de poder reinterpretar el análisis realizado en [5] utilizando herramientas propias del análisis no estándar. Con este enfoque, las matrices de rotación infinitesimales tendrían como elementos números infinitesimales —lo cual no presenta un problema, dado que ${}^*\mathbb{R}$ es un cuerpo—, y bastaría con definir que dos vectores están infinitesimalmente cerca si la diferencia entre ellos tiene todas sus componentes infinitesimales.

Para que la monografía no se extienda más de lo debido, se optó por no desarrollar en mayor profundidad esta aplicación, puesto que se debería realizar una introducción a conceptos propios de mecánica que, a diferencia del desarrollo de la Sección 4.2, requieren conocimientos más técnicos de física.

5. Epílogo

*“Todo tiene su final.
Nada dura para siempre.”*

— Héctor Lavoe y Willie Colón,
Todo tiene su final

Antes de finalizar este trabajo, me gustaría agradecer al lector que ha llegado hasta aquí. Si bien es cierto que, en ocasiones, el contenido teórico pudo haber resultado arduo de incorporar en una primera lectura, espero que haya despertado su curiosidad y dejado algunas inquietudes que merezcan ser exploradas.

A lo largo de esta monografía se pudo ver la construcción, el desarrollo y algunas aplicaciones del análisis no estándar. Cada uno de estos aspectos puede profundizarse aún más, en función de los conocimientos y los gustos que tenga el lector. Sin embargo, quizás lo más valioso de este texto haya sido el cambio de perspectiva que presenta.

Me gustaría destacar el hecho de que, si bien la construcción de los números hiperreales puede resultar muy “técnica”, esto no es un impedimento para su uso. La importancia del análisis no estándar radica en conocer las propiedades de los números hiperreales, y cómo utilizar las mismas para obtener resultados. Los formalismos necesarios para construir el conjunto numérico con el cual se trabaja son simplemente eso, formalismos necesarios, pues nadie piensa en los números reales como clases de equivalencia de sucesiones de Cauchy de números racionales a la hora de calcular una derivada.

El análisis no estándar nos brinda herramientas para mirar con nuevos ojos estructuras conocidas, y nos invita a reflexionar sobre conceptos ya incorporados en nuestra forma actual de estudiar y entender la matemática. Porque a veces, cuando una idea surge de una intuición tan profunda, ni el tiempo ni la ausencia son suficientes para matarla.

Bibliografía

- [1] L. O. Arkeryd, N. J. Cutland, and C. Ward Henson, editors. *Nonstandard Analysis: Theory and Applications*, volume 493 of *NATO ASI Series C: Mathematical and Physical Sciences*. Springer, Dordrecht, 1997.
- [2] A. L. Cauchy. *Cours d'analyse de l'École royale polytechnique*. Imprimerie Royale, chez Debure frères, Paris, 1821. <https://gallica.bnf.fr/ark:/12148/btv1b8626657t>.
- [3] C. H. Edwards. *The Historical Development of the Calculus*. Springer, New York, NY, 1 edition, 1979.
- [4] R. Goldblatt. *Lectures on the Hyperreals: An Introduction to Nonstandard Analysis*, volume 188 of *Graduate Texts in Mathematics*. Springer, New York, 1998.
- [5] H. Goldstein, C. P. Poole, and J. L. Safko. *Classical Mechanics*. Addison-Wesley, San Francisco, 3rd edition, 2002.
- [6] H. J. Keisler. *Elementary Calculus: An Infinitesimal Approach*. Dover Publications, 2013. Reprint of the 1976 University of Wisconsin edition.
- [7] A. M. Robert. *Nonstandard Analysis*. Wiley-Interscience, New York, 1988.
- [8] A. Robinson. *Non-Standard Analysis*. North-Holland (reimpresión Princeton University Press), Amsterdam / Princeton, 1966. Revised ed. disponible.
- [9] R. A. Serway and J. W. Jewett Jr. *Physics for Scientists and Engineers with Modern Physics*. Cengage Learning, Belmont, CA, 7 edition, 2007.