

Indice

Presentazione	xiii
Introduzione	xv
1 Introduzione: $\mathbb{Q}$ non basta	1
1.1 Introduzione	1
1.2 La radice quadrata	2
1.2.1 La radice quadrata di 2	2
1.2.2 Le radici quadrate di 3, 5, ..., 17	3
1.2.3 La radice quadrata degli interi non quadrati	6
1.3 L'estremo superiore	7
2 Definizione assiomatica di $\mathbb{R}$	11
2.1 Introduzione	11
2.2 Campi ordinati	13
2.3 Campi ordinati archimedei	17
2.4 Campi ordinati completi	20
2.5 Teorema di Hilbert	21
2.5.1 Omomorfismi ordinati	21
2.5.2 "Unicità"	26
2.5.3 Definizione assiomatica di $\mathbb{R}$	26
2.6 L' <i>Axiom der Vollständigkeit</i>	27
2.6.1 Estensioni di campi	28
2.6.2 Estensioni di campi ordinati	29
2.6.3 La <i>Vollständigkeit</i> equivale alla completezza	31
2.6.4 "Esistenza": un'idea per un modello di $\mathbb{R}$	33
2.6.5 Una precisazione terminologica	35

2.7	Proprietà di Dedekind e di Cantor	36
2.8	Commenti	41
3	La costruzione di Méray–Cantor	43
3.1	Introduzione	44
3.2	La costruzione di Méray–Cantor	45
3.2.1	Definizione dei numeri reali	47
3.2.2	Immersione dei razionali	50
3.2.3	La struttura d'ordine	50
3.2.4	Il teorema di completezza <i>à la</i> Méray–Cantor . .	53
3.3	Sulla definizione di successione di Cauchy	59
3.3.1	La definizione di Kline	59
3.3.2	La definizione di Bolzano	62
3.3.3	Una definizione basata sulla finitezza	63
4	La costruzione di Dedekind	65
4.1	Introduzione	65
4.2	Le sezioni di Dedekind	66
4.2.1	Definizione dei numeri reali	66
4.2.2	Immersione dei razionali	68
4.2.3	La struttura d'ordine	71
4.2.4	Il teorema fondamentale della costruzione	72
4.2.5	Il teorema di completezza <i>à la</i> Dedekind	74
4.2.6	La struttura algebrica	75
4.3	Un percorso alternativo	82
4.4	Commenti	84
5	La completezza	87
5.1	Introduzione	87
5.2	Il circolo della completezza	87
5.3	Inizia l'analisi matematica	90
6	Gli allineamenti di cifre	93
6.1	Introduzione	93
6.2	Gli allineamenti binari	94
6.2.1	La rappresentazione binaria	95
6.2.2	Implementazione in R	99
6.2.3	Implementazione in <i>Mathematica</i>	103

6.3	Costruzione di $\mathbb{R}$ tramite allineamenti	104
6.4	Costruzione di $\mathbb{R}$ tramite classi contigue	105
7	Gli iperreali	107
7.1	Introduzione	107
7.2	Costruzione degli iperreali	109
7.2.1	Filtri e ultrafiltri	109
7.2.2	Il Lemma di Zorn	111
7.2.3	Ultrafiltri non principali	113
7.2.4	Gli iperreali	113
7.2.5	“Unicità”	115
7.2.6	Immersione dei reali	116
7.2.7	La struttura d’ordine	116
7.2.8	Esistenza degli infinitesimi	118
7.2.9	Confronto fra infinitesimi	119
7.2.10	Esistenza degli infiniti	120
7.2.11	Teorema della parte standard	120
7.3	Cenni allo sviluppo dell’analisi non-standard	122
8	Numeri reali e	127
8.1	Introduzione	127
8.2	Numeri reali e cardinalità	128
8.2.1	Insiemi numerabili	130
8.2.2	Il procedimento diagonale di Cantor	131
8.2.3	Non-numerabilità di $[0, 1)$	131
8.2.4	Numeri razionali, algebrici, calcolabili	133
8.2.5		134
8.2.6	Un esempio di numero reale non-calcolabile	135
8.2.7	Teorema di Turing	135
8.2.8	Numeri definibili e Paradosso di Richard	136
8.2.9	La “definizione” di successione	138
8.3	Numeri reali e misura	139
8.3.1	Insiemi trascurabili	139
8.3.2	L’insieme di Cantor	141
8.3.3	Numeri normali	142
8.4	Numeri reali e scienze cognitive, <i>di Michela Maschietto</i>	143
8.4.1	La scienza cognitiva della matematica	144

8.4.2	Invarianti e fondamenti della matematica	149
A	Postfazione	153
A.1	Un modello “didattico” di $\mathbb{R}$	154
A.1.1	Segmenti di $\mathbb{Q}^+$	154
A.1.2	Operazioni ed ordine fra segmenti di $\mathbb{Q}^+$	155
A.1.3	Proprietà delle operazioni e dell’ordine	156
A.1.4	Completezza	157
A.1.5	Conclusione della costruzione	157
	Bibliografia	159
	Indice analitico	167
	Glossario essenziale	169

Elenco delle figure

1	Interdipendenza dei capitoli	xxi
1.1	Grafico di $f(x) = -2x + \frac{x^3}{3}$	8
2.1	David Hilbert ed Archimede	12
3.1	Hugues Méray e Georg Cantor	43
3.2	Teorema fondamentale di Méray–Cantor	56
4.1	Richard Dedekind ed Euclide	65
5.1	Circolo della completezza	91

Elenco delle tabelle

1.1	Tabulazione di x^2 in $\mathbb{Z}_a$	4
2.1	Legami fra le proprietà di Cantor, di Archimede, di Dedekind e la completezza nei campi ordinati	39

Presentazione

I numeri reali, intesi come allineamenti decimali (o binari) con “code” arbitrariamente lunghe e potenzialmente illimitate compaiono in ambito scolastico fin dal primo ciclo (a livello della ex scuola media), diventano poi vieppiù familiari con l’uso degli strumenti di calcolo elettronico, e si rivelano indispensabili per quantificare le misure delle grandezze in tutte le discipline scientifiche.

Quando poi, nel prosieguo degli studi secondari e universitari, si passa all’introduzione del campo dei numeri reali e allo studio delle sue proprietà strutturali, i principali punti di riferimento vengono ad essere (oltre alla teoria delle grandezze di Eudosso, che ragionava in termini geometrici) le costruzioni di Méray–Cantor e di Dedekind, e la caratterizzazione assiomatica di Hilbert.

Purtroppo nella prassi corrente ci si ferma spesso a questo punto, tralasciando di inquadrare adeguatamente il metodo degli allineamenti decimali (o binari) nell’ambito delle costruzioni teoriche or ora citate. Col risultato che in molti discenti rimane la sgradevole sensazione di avere a che fare con due mondi separati.

Queste riflessioni mi sono state suggerite (insieme a molte altre che per motivi di brevità non posso esplicitare in questa sede) dalla lettura dello stimolante testo di Carla Fiori e Sergio Invernizzi, dove invece il suddetto inquadramento è chiaramente delineato e arricchito con esempi sorprendenti quali ad esempio le nozioni di numeri reali “non calcolabili” e numeri reali “non definibili”.

Sono certo che il libro sarà apprezzato dai lettori e risulterà particolarmente utile per un arricchimento culturale e professionale dei docenti in servizio e per la formazione dei futuri docenti di matematica.

Vinicio Villani

Introduzione

Mi sembra importante che gli insegnanti siano consapevoli dell'esistenza di una pluralità di approcci diversi [alla costruzione dei numeri reali] e dei loro mutui legami.

VINICIO VILLANI,
Cominciamo da Zero

La “scoperta” dei numeri non razionali o più precisamente delle grandezze non commensurabili risale a più di 2000 anni fa, ma solo nel XIX secolo, con lo sviluppo dell'analisi matematica, nasce l'esigenza di una definizione rigorosa di “numero reale” ed emergono varie problematiche di carattere logico e filosofico tutt'altro che elementari e non ancora del tutto risolte. Ancora oggi il sistema dei numeri reali è oggetto di studi ed approfondimenti, non solo in ambito matematico ed epistemologico, ma anche in ambiti quali ad esempio quello delle scienze cognitive e delle neuroscienze.

Questo lavoro presenta unitariamente i principali metodi per definire i numeri reali, ripropone le dimostrazioni classiche, ma presenta anche alcuni contributi che si discostano dalle usuali trattazioni, sia in alcune dimostrazioni che nelle parti comparative. L'intenzione degli autori è di fornire uno strumento utile per una corretta trattazione e per una rivisitazione comparativa e critica della teoria dei numeri reali, nonché di fornire spunti di riflessione.

Il testo è rivolto a tutti coloro che desiderano “saperne di più” sul sistema dei numeri reali e non solo agli studiosi di matematica. Per questo l'esposizione, pur essendo rigorosa e dettagliata, è stata mantenuta

entro livelli di linguaggio e di complessità accessibili anche ad un lettore “non matematico”. Per le parti più specialistiche, al lettore interessato vengono fornite indicazioni e riferimenti bibliografici esplicativi e utili per gli approfondimenti.

L’auspicio è comunque che questo lavoro sia di aiuto, in primo luogo, a coloro che insegnano materie scientifiche o che si apprestano a farlo. Di norma, in ambito scolastico, i numeri reali vengono trattati solo come strumenti di calcolo, ma questo, come ha osservato Laure Barthel [5], induce un’idea di numero reale molto lontana da una definizione formale e dalle problematiche ad essa legate. D’altra parte, recenti ricerche mostrano che il cervello umano non ha un modello intuitivo per i numeri reali, a differenza di ciò che accade per i numeri naturali. Una comprensione distorta della natura dei numeri reali ha conseguenze più ampie della semplice mancanza di comprensione di una corretta formalizzazione. Per esempio, ingegneri ed informatici hanno bisogno di una buona conoscenza del legame fra un numero reale e la sua rappresentazione decimale (o binaria), per poter cogliere le problematiche algoritmiche relative alla rappresentazione dei numeri reali nei computer. Si aggiunga che quello dei numeri reali è un argomento che può dare ottime possibilità per lavori interdisciplinari.

La spinta più forte a scrivere questo libro è venuta da problemi e considerazioni emersi, anno dopo anno, durante l’attività didattica degli autori maturata in corsi universitari di diverse tipologie.

Punto di partenza è stata la constatazione che l’aggettivo *completo* compare tipicamente nell’insegnamento universitario in ben tre situazioni diverse:

- nella storia e nei fondamenti della matematica, nell’assioma di completezza di Hilbert (*Axiom der Vollständigkeit*);
- in analisi matematica, nella nozione di spazio normato completo;
- in algebra, nella nozione di campo ordinato completo.

Le tre nozioni di completezza richiamate hanno le radici nel teorema di isomorfismo di Hilbert, nel modello dei numeri reali di Cantor e nel modello dei numeri reali di Dedekind, e danno rispettivamente origine ai Capitoli 2, 3 e 4 di questo libro. In particolare, questi tre capitoli intendono “solidificare” la difficile costruzione dei reali, cercando di col-

legare fra loro le tre colonne (Hilbert, Cantor, Dedekind) con una buona architrave.

Un secondo fatto è stata la considerazione che come punto di partenza dell'analisi viene tradizionalmente scelto il teorema di esistenza dell'estremo superiore. In tal modo si assume implicitamente il modello di Dedekind, e si ottiene il risultato didattico di rinviare di qualche lezione la delicata dimostrazione del criterio di convergenza di Cauchy. Tale prassi didattica ed eventuali alternative sono discusse nel Capitolo 5.

In tutto il testo abbiamo deliberatamente taciuto dei problemi strettamente logici e fondazionali connessi sia alle costruzioni dei numeri reali che alla successiva costruzione dell'analisi e di gran parte della matematica. Questo se non altro per fedeltà storica, dato che la cosiddetta aritmetizzazione dell'analisi – ovvero la fondazione dei reali sui razionali – ha sostanzialmente preceduto la crisi dei fondamenti, sviluppata nel primo trentennio del XX secolo. L'omissione è quindi dovuta alla convinzione che, prima di studiare la crisi dei fondamenti, conviene averne ben chiare le cause, fra le quali certamente si ascrivono le costruzioni dei reali sui razionali conseguite negli anni '70 del XIX secolo.

Veniamo ad una descrizione più dettagliata dei contenuti, indicando, capitolo per capitolo, quali sono i punti principali per i quali non abbiamo trovato un riscontro specifico nella letteratura esistente. La Figura 1 illustra lo schema di interdipendenza dei capitoli.

Il Capitolo 1, introduttivo, richiama ed analizza la dimostrazione classica di irrazionalità di $\sqrt{2}$. Per lo studio dei casi “platonici” (la radice di a , con $3 \leq a \leq 17$) viene proposta una dimostrazione che *potrebbe* essere stata alla portata dei matematici greci, e che, avendo complessità di calcolo crescente con a , potrebbe spiegare perché ci si fermò a 17. Per il caso generale, riportiamo la brevissima dimostrazione di Richard Beigel. Segnaliamo infine il Teorema 1.2.2, che rappresenta un interessante esercizio di aritmetica modulare.

Il Capitolo 2 presenta l'introduzione assiomatica di $\mathbb{R}$ come campo ordinato *completo*, cioè come un campo ordinato in cui ogni sottoinsieme non-vuoto e superiormente limitato ammette l'estremo superiore. Viene dimostrata, in modo classico, l'“unicità” di una tale struttura, cioè provato che due campi ordinati completi sono ordinatamente isomorfi (Teorema 2.5.7). Viene poi presentata la definizione storica di $\mathbb{R}$ data da

Hilbert, basata sull'*Axiom der Vollständigkeit*, e se ne dimostra l'equivalenza con quella assiomatica attuale (Teorema 2.6.12). Questa parte, per la quale gli autori non hanno trovato riferimenti in letteratura, mostra in particolare che la Proprietà di Archimede (ossia l'illimitatezza dei numeri naturali nei reali) svolge un ruolo essenziale nella definizione di $\mathbb{R}$. Infine, per le implicazioni didattiche, vengono discusse le definizioni assiomatiche di $\mathbb{R}$ basate sulle *proprietà di separazione*: la Proprietà di Cantor (esistenza di elemento separatore fra classi contigue) e la Proprietà di Dedekind (esistenza di elemento separatore fra classi separate), con le sue diverse varianti presenti nei libri di testo.

Il Capitolo 3 presenta integralmente la costruzione di $\mathbb{R}$ di Méray–Cantor basata sulle successioni di Cauchy, fino a provare la *completezza metrica* di $\mathbb{R}$ (ogni successione reale di Cauchy converge). La costruzione può essere utilizzata anche in un corso di algebra come esempio non banale di applicazione del teorema sul quoziente di un anello commutativo unitario su un ideale massimale. La sezione conclusiva del Capitolo 3 studia il problema didattico (e logico) della definizione classica di successione di Cauchy (con quattro quantificatori), proponendo la definizione alternativa di Bolzano (con tre quantificatori) e la definizione basata sulla finitezza (con un solo quantificatore) utilizzata in alcuni verificatori automatici. La Sezione 3.3 è dovuta ad una collaborazione con Antonella Montone; in essa si discutono le problematiche didattiche del concetto di successione di Cauchy, segnalando anche un'errata traduzione o interpretazione, presente in taluni autori, della definizione originale (di successione di Cauchy) data da Cantor.

Il Capitolo 4 presenta la costruzione di Dedekind basata sulle sezioni di $\mathbb{Q}$. Le noiose verifiche formali richieste da questa costruzione vengono presentate mettendo in particolare evidenza i punti dove si usa la archimedeicità di $\mathbb{Q}$. Per il modello di $\mathbb{R}$ di Dedekind viene ovviamente provata la *completezza* nel senso dei corpi ordinati (esistenza dell'estremo superiore).

È bene osservare che, storicamente, i risultati del Capitolo 2 seguono di circa trent'anni quelli dei Capitoli 3 e 4, dei quali costituiscono, in un certo senso, la “sistemazione” teorica. Nel testo, per comodità, abbiamo preferito ribaltare quest'ordine cronologico, ed abbiamo anteposto la definizione assiomatica in modo da aver disponibile il quadro teorico di riferimento al momento della costruzione dei due modelli classici di $\mathbb{R}$.

Il Capitolo 5 mostra come il Teorema di Cauchy (ogni successione reale di Cauchy converge), “coronamento” della costruzione di Méray–Cantor, ed il Teorema di Bolzano–Weierstrass (essenzialmente la compattezza sequenziale dei chiusi limitati), a sua volta “coronamento” della costruzione di Dedekind, siano dimostrabili l’uno dall’altro. Il percorso didattico tradizionale italiano, che di norma utilizza un’introduzione assiomatica, preferisce dimostrare il Teorema di Cauchy usando il Teorema di Bolzano–Weierstrass. A conoscenza degli autori non appare sperimentato l’equivalente percorso inverso.

Il Capitolo 6 discute l’introduzione di $\mathbb{R}$ con gli allineamenti di cifre decimali. Per semplicità vengono usati gli *allineamenti binari*. Un allineamento di cifre viene interpretato non tanto come lista dei coefficienti di una serie di potenze, ma come *codice* per la descrizione di una sezione di Dedekind e quindi di un numero reale. Come esercizi, vengono proposte implementazioni in linguaggio R di alcuni sviluppi binari (ad esempio della $\sqrt{2}$ a partire dalla sua definizione come sezione). La conclusione principale del Capitolo 6 è che la costruzione di $\mathbb{R}$ con gli allineamenti decimali è nient’altro che un caso particolare della costruzione di Dedekind, della quale conseguentemente conserva i difetti: le ambiguità del tipo $0.99999\dots = 1.00000\dots$ nella rappresentazione decimale dei reali hanno la loro ineliminabile radice nell’ambiguità della rappresentazione di una sezione di $\mathbb{Q}$ di prima specie, nella quale il numero r “dove si taglia” può indifferentemente essere messo nella classe inferiore o nella classe superiore. Questa “spiegazione” dell’uguaglianza $0.99999\dots = 1.00000\dots$ ci appare molto più convincente della mera constatazione che $\sum_{k=1}^{\infty} 9/10^k = 1$.

Il Capitolo 7 è essenzialmente un *bonus* aggiunto al Capitolo 3. In esso mostriamo infatti come la costruzione di Méray–Cantor possa essere ripetuta passo a passo *mutatis mutandis* per costruire i numeri iperreali a partire dai reali. Inoltre il tema degli iperreali fornisce un esempio importante di corpo ordinato non–archimedeo. La costruzione degli iperreali viene preceduta da brevi richiami sugli ultrafiltri. Il capitolo si conclude con il Teorema della Parte Standard e con la definizione non–standard di derivata di una funzione $\mathbb{R} \rightarrow \mathbb{R}$.

Il Capitolo 8 contiene vari spunti di riflessione sul concetto di numero reale. Gli argomenti presentati mostrano la grande delicatezza, anche in termini filosofici, della nozione di numero reale. Sono per prime

esaminare alcune questioni legate alla non-numerabilità di $\mathbb{R}$. Vengono richiamati i sottoinsiemi numerabili “classici”, quello dei numeri razionali e quello dei numeri algebrici, e viene introdotto, in modo informale, il sottoinsieme numerabile dei numeri calcolabili. Si mostra che i numeri dei corrispondenti insiemi complementari (irrazionali, trascendenti, non-calcolabili) costituiscono, da diversi punti di vista, la “stragrande maggioranza” dei numeri reali, con la sconcertante conclusione che “quasi tutti” i numeri reali non possono essere né mai saranno “calcolati” dall’uomo. Questa discussione sui numeri calcolabili e non-calcolabili costituisce un importante argomento logico-informatico, che spesso è assente dalle trattazioni più strettamente matematiche sui numeri reali. Viene anche proposta una “soluzione” del Paradosso di Richard (legato ai numeri “definibili” e “non-definibili”) ispirata direttamente al Teorema di Indecidibilità di Turing. Questa parte può anche essere di stimolo per attività interdisciplinari (con l’informatica, con la logica, con la filosofia). Il capitolo è infine concluso da un contributo di Michela Maschietto relativo alla genesi del concetto di numero reale secondo alcuni recenti studi che utilizzano ed applicano le tecniche delle scienze cognitive e delle neuroscienze.¹ Dal punto di vista dello sviluppo concettuale, i numeri reali hanno una struttura assolutamente diversa da quella della numerosità, legata al contare ed ai numeri naturali, in quanto non avrebbero, a differenza dei numeri naturali, un analogo “neuronal” nel cervello umano. Per la comprensione dei numeri reali vi sarebbe quindi, in aggiunta agli ostacoli *epistemologici* noti, un ostacolo *biologico* non eliminabile.

¹ In particolare, il contributo fa riferimento alla teoria dell’*embodied cognition* (o cognizione incarnata), secondo la quale le idee e le altre attività della mente umana dipendono strettamente dalle varie attività del corpo, quali la percezione del mondo esterno, il movimento, e le altre interazioni con l’ambiente. Pur non generalmente accettata dagli psicologi, l’*embodied cognition* ha destato un notevole interesse fra i matematici, dopo che due dei suoi “fondatori”, George P. Lakoff e Rafael E. Núñez [35], hanno proposto nel suo ambito una teoria sull’origine dei concetti della matematica.

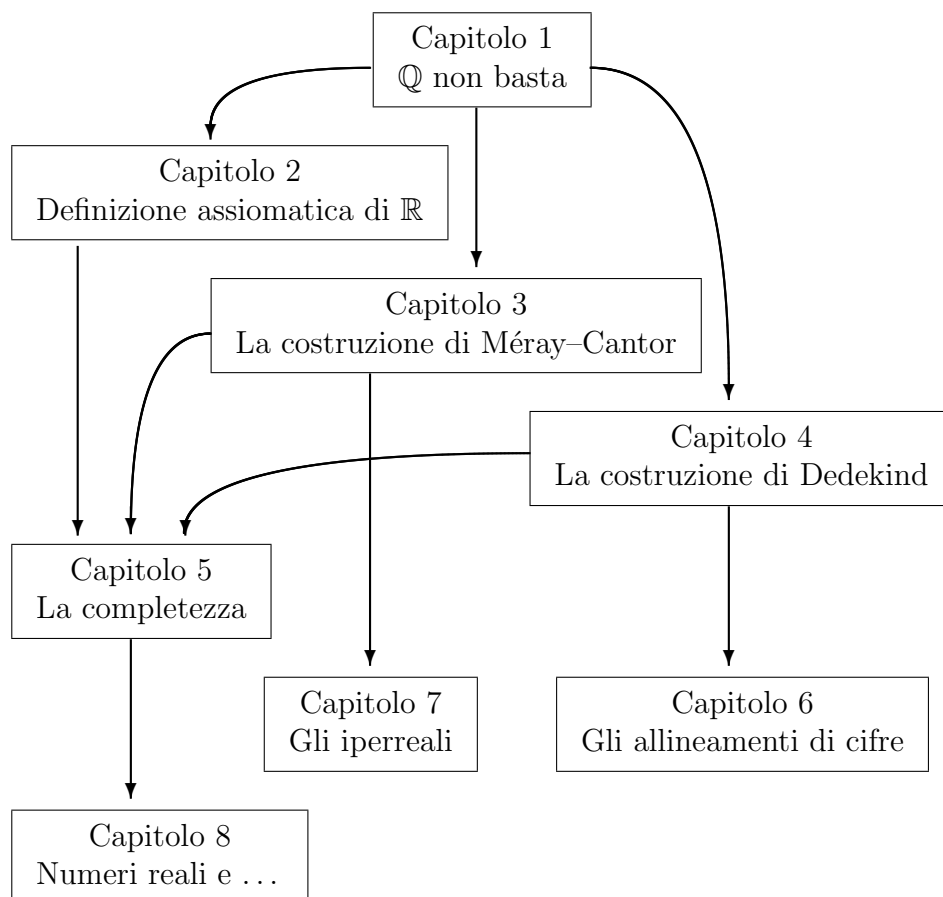


Figura 1: Schema dell’interdipendenza dei capitoli. Si noti che i capitoli centrali 2, 3 e 4 possono essere letti in modo indipendente. Il Capitolo 8 è essenzialmente indipendente dal resto del testo e può essere letto anche senza una conoscenza precisa della struttura di $\mathbb{R}$, anzi, può servire da “stimolo” per un successivo approfondimento; nello schema viene fatto dipendere dal Capitolo 5 unicamente per l’uso del Teorema di Heine–Borel nella dimostrazione della Proposizione 8.3.3.

RINGRAZIAMENTI

Gli autori desiderano ringraziare Giorgio Bagni, Giulio Cesare Barozzi, Bruno D'Amore, Luigi Maierù, Antonella Montone, Eugenio Omodeo, Andrea Sgarro, Walter Spangher e Vinicio Villani per alcuni utili colloqui e suggerimenti. Ringraziamo anche, per il loro entusiasmo e le loro osservazioni, gli studenti della laurea specialistica in matematica dell'Università di Trieste, dove questo libro è stato sperimentato come testo per l'insegnamento di *Matematiche Complementari I* negli aa.aa. 2006–2007 e 2007–2008, e gli studenti della laurea specialistica in matematica dell'Università di Modena e Reggio Emilia dove i Capitoli 1, 2, 3 e 8 sono stati discussi in seminari di approfondimento per il corso di *Strutture Algebriche*.

Il lavoro è stato finanziato dai PRIN “Strutture geometriche, combinatoria e loro applicazioni” (C. Fiori) e “Sviluppo su grande scala di dimostrazioni certificate” (S. Invernizzi).

Modena–Trieste, ottobre 2008

Capitolo 1

Introduzione: $\mathbb{Q}$ non basta

“E allora, zio,” dissi, in tono bellicoso, “fra un anno prenderò la laurea e mi sto già preparando per l’ammissione alla scuola di perfezionamento. La tua manovra è fallita. Ti piaccia o no, diventerò un matematico.”

APOSTOLOS DOXIADIS,
Zio Petros e la Congettura di Goldbach

1.1 Introduzione

L’insieme $\mathbb{Q}$ dei numeri razionali non basta per le necessità dell’analisi matematica. Questa affermazione si trova tipicamente all’inizio di ogni testo di analisi matematica (cf. ad es. [22], p. 14). La giustificazione classica è il fatto che non esiste alcun numero razionale a tale che $a^2 = 2$. In questo capitolo introduttivo, richiamiamo ed analizziamo la dimostrazione classica di irrazionalità di $\sqrt{2}$. La dimostrazione pitagorica di questo teorema non è nota, ma se si basava, come appare in Aristotele (cf. [16], p. 225), sulla differenza fra numeri pari e numeri dispari, poteva essere quella presentata nel Teorema 1.2.1 ([22], Proposizione 3.1). Per lo studio dei casi “platonici” (la radice di a , con $3 \leq a \leq 17$) viene proposta una dimostrazione che *potrebbe* essere stata alla portata dei matematici greci, e che, avendo complessità di calcolo crescente con a , potrebbe spiegare perché ci si fermò a 17. Per il caso generale, riportiamo la brevissima

dimostrazione di Richard Beigel. Segnaliamo infine il Teorema 1.2.2, che rappresenta un interessante esercizio di aritmetica modulare.

1.2 La radice quadrata

1.2.1 La radice quadrata di 2

Teorema 1.2.1 *Non esistono interi positivi n, m tali che $n^2 = 2 \cdot m^2$.*

DIMOSTRAZIONE. Supponiamo che esistano interi positivi n, m primi tra loro (privi di divisori comuni $\neq 1$) tali che $n^2 = 2 \cdot m^2$. Ripartiamo i numeri interi positivi in due classi, pari (del tipo $2k$, $k \geq 1$) e dispari (del tipo $2k + 1$, $k \geq 1$). Si ha

$$\begin{cases} \text{il quadrato di un numero pari è pari,} \\ \text{il quadrato di un numero dispari è dispari,} \end{cases} \quad (1.1)$$

ossia un numero intero positivo ha quadrato pari se e solo se è pari. Ora da $n^2 = 2 \cdot m^2$ segue che n^2 è pari, e quindi n è pari; inoltre, poiché n ed m sono privi di divisori comuni, si ha che m è dispari, e quindi m^2 è dispari. Pertanto $n^2 = 2 \cdot m^2$ è assurdo, perché $2 \cdot m^2$ è divisibile per 2 ma non per 4, mentre n^2 è divisibile per 4. $\square$

Cerchiamo ora di complicare (almeno formalmente) la stessa dimostrazione.¹

DIMOSTRAZIONE. Supponiamo che esistano interi positivi n, m primi tra loro (privi di divisori comuni $\neq 1$) tali che $n^2 = 2 \cdot m^2$. Ripartiamo l'insieme $\mathbb{Z}$ dei numeri interi in due classi modulo 2, la classe $[0]$ degli interi pari (del tipo $2k$, $k \in \mathbb{Z}$) e la classe $[1]$ degli interi dispari (del tipo $2k + 1$, $k \in \mathbb{Z}$). In $\mathbb{Z}_2$ si ha

$$\begin{cases} [0]^2 = [0] \\ [1]^2 = [1] \end{cases} \quad (1.2)$$

ossia $[x]^2 = [0]$ se e solo se $[x] = [0]$. Ora da $n^2 = 2 \cdot m^2$ segue $[n^2] = [n]^2 = [0]$, e quindi $[n] = [0]$. Dato che n ed m sono privi di divisori comuni,

¹ Useremo i concetti di *anello* e di *classi di resto modulo a* (dette anche *classi di congruenza modulo a*), per i quali si veda [26], p. 130 e p. 23.

si ha che $[m] = [1]$, e quindi $[m^2] = [1] \neq [0]$. Pertanto $n^2 = 2 \cdot m^2$ è assurdo, perché $2 \cdot m^2$ è divisibile per 2 ma non per 2^2 , mentre n^2 è divisibile per 2^2 . $\square$

1.2.2 Le radici quadrate di 3, 5, ..., 17

La seconda dimostrazione del Teorema 1.2.1 sopra riportata pare una inutile complicazione della dimostrazione classica, ma ci suggerisce il seguente:

Teorema 1.2.2 *Sia $a \geq 2$ un intero tale che nell'anello $\mathbb{Z}_a$ delle classi di resto modulo a l'equazione $[x]^2 = [0]$ abbia solo la soluzione $[x] = [0]$. Allora non esistono interi positivi n, m tali che $n^2 = a \cdot m^2$.*

DIMOSTRAZIONE. Supponiamo che esistano interi positivi n, m primi tra loro (privi di divisori comuni $\neq 1$) tali che $n^2 = a \cdot m^2$. Ora da $n^2 = a \cdot m^2$ segue $[n^2] = [a] = [0]$, ma $[n^2] = [n]^2$, e quindi $[n] = [0]$, ossia $n = ha$ (con h intero) e perciò n^2 è divisibile per a^2 . Poiché n ed m sono privi di divisori comuni, si ha che $[m] \neq [0]$, da cui $[m^2] \neq [0]$ e pertanto $a \cdot m^2$ è divisibile per a ma non per a^2 . Nell'ipotesi $n^2 = a \cdot m^2$ si giunge quindi ad un assurdo. $\square$

In particolare il Teorema 1.2.2 implica che i numeri primi p non hanno radice quadrata razionale: infatti, se p è un numero primo, $\mathbb{Z}_p$ è un corpo, e quindi $[x]^2 = [0]$ solo se $[x] = [0]$.

Come applicazione, per esercizio, consideriamo gli interi a , $3 \leq a \leq 17$, come nel famosissimo passo del *Teeteto* di Platone [48]:

...

Socrate: E quale questa maniera, Teeteto?

Teeteto: Teodoro disegnava un qualcosa per noi intorno alle potenze, per esempio su quella di tre piedi e su quella di cinque, mostrando che per lunghezza queste potenze non si possono misurare con la lunghezza di un piede. E così scegliendo una per ciascuna le potenze giunse fino alla diciassettesima e in questa si trattenne. ...²

² Σωκράτης· τὸ ποῖον δὴ, ὦ Θεαίτητε; Θεαίτητος· περὶ δυνάμεων τι ἡμῖν Θεόδωρος ὁδε ἔγραψε, τῆς τε τρήποδος πέρι καὶ πεντέποδος [ἀποφαίνων] ὅτι μήκει οὐ σύμμετροι τῇ ποδιαίᾳ, καὶ οὕτω κατὰ μίαν ἐκάστην προαιρούμενος μέχρι τῆς ἐπτακαίδεκάποδος· ἐν δὲ ταύτῃ πως ἐνέσχετο...

Utilizziamo un software, ad esempio il pacchetto R [49],³ per tabulare nell'anello $\mathbb{Z}_a$ i valori $x^2 \bmod (a)$, al variare di $x \in \mathbb{Z}_a$, ossia per $x = 0, 1, \dots, a - 1$. La tabulazione è riportata nella Tabella 1.1, in cui ogni riga corrisponde ad un valore di a per $1 \leq a \leq 17$.

[1]	0																
[2]	0	1															
[3]	0	1	1														
[4]	0	1	0	1													
[5]	0	1	4	4	1												
[6]	0	1	4	3	4	1											
[7]	0	1	4	2	2	4	1										
[8]	0	1	4	1	0	1	4	1									
[9]	0	1	4	0	7	7	0	4	1								
[10]	0	1	4	9	6	5	6	9	4	1							
[11]	0	1	4	9	5	3	3	5	9	4	1						
[12]	0	1	4	9	4	1	0	1	4	9	4	1					
[13]	0	1	4	9	3	12	10	10	12	3	9	4	1				
[14]	0	1	4	9	2	11	8	7	8	11	2	9	4	1			
[15]	0	1	4	9	1	10	6	4	4	6	10	1	9	4	1		
[16]	0	1	4	9	0	9	4	1	0	1	4	9	0	9	4	1	
[17]	0	1	4	9	16	8	2	15	13	13	15	2	8	6	9	4	1

Tabella 1.1: Tabulazione di x^2 in $\mathbb{Z}_a$, ossia dei valori $x^2 \bmod (a)$, per $a = 1, 2, \dots, 17$ (righe) e $x = 0, 1, \dots, a - 1$ (colonne). La tabulazione è stata ottenuta in R con il semplice comando `for (a in (1:17)) {print(mod((0:(a-1))^2,a))}`, essendo la funzione `mod` (modulo) definita da `mod <- function(n,m){n-(n%/%m)*m}`.

Per esempio in $\mathbb{Z}_6$, $\mathbb{Z}_{10}$, $\mathbb{Z}_{14}$, $\mathbb{Z}_{15}$, si ha $x^2 \neq 0$ se $x \neq 0$, per cui i numeri 6, 10, 14 e 15, non hanno radice quadrata razionale. Idem per 2, 3, 5, 7, 11, 13, 17, che sono primi e per i quali la tabulazione conferma che $x^2 = 0$ implica $x = 0$. Naturalmente dobbiamo “saltare” i *quadrati*⁴

³ R è un sistema “open source” disponibile nei termini di una GPL; è gratuito e disponibile per molte piattaforme UNIX, Windows e MacOS. Anche se pensato originariamente per l’analisi statistica, la simulazione e la rappresentazione grafica dei dati, R ha capacità didattiche estremamente interessanti.

⁴Un intero positivo a è detto un *quadrato* se esiste un intero positivo b tale che $a = b^2$.

4, 9, 16, che probabilmente – anche se non detto esplicitamente – erano saltati anche dal nostro Teodoro. Restano $a = 8$, $a = 12$, per i quali in $\mathbb{Z}_a$ ci sono soluzioni non nulle di $x^2 = 0$. Ciò non stupisce perché sia 8 che 12 sono del tipo $a = bc^2$, ossia multipli di quadrati (in particolare $8 = 2 \cdot 2^2$, $12 = 3 \cdot 2^2$), e quindi $x = bc$ è una soluzione di $x^2 = 0$ (infatti $x^2 = b^2c^2 = b \cdot bc^2 = ba$ è multiplo di a).

Per dimostrare che $a = 8$ e $a = 12$ non hanno radice quadrata razionale r , possiamo procedere in due modi. Primo modo: se r fosse un razionale tale che $r^2 = bc^2$, allora r/c sarebbe un razionale tale che $(r/c)^2 = b$, e quindi la non-razionalità della radice di 8 e di 12 viene ricondotta rispettivamente a quella delle radici di 2 e di 3.

Secondo modo: si tratta sostanzialmente di un esercizio di aritmetica modulare.

Sia $a = 8$. Supponiamo che esistano interi positivi n, m primi tra loro tali che $n^2 = 8 \cdot m^2$. Ora da $n^2 = 8 \cdot m^2$ segue che in $\mathbb{Z}_8$ si ha $[n^2] = [n]^2 = [0]$, e quindi (vedi Tabella 1.1) possiamo avere due casi: $[n] = [0]$, oppure $[n] = [4]$.

- Caso $[n] = [0]$.
 - Se $[m^2] \neq [0]$: si procede con la solita prova; $8 \cdot m^2$ è divisibile per 8 ma non per 8^2 , mentre n^2 è divisibile per 8^2 .
 - Se $[m^2] = [0]$: sempre dalla Tabella 1.1 ci sono due possibilità:
 - * $[m] = [0]$: si esclude perché n, m sono primi fra loro.
 - * $[m] = [4]$: ossia $n = 8k$, $m = 8h + 4$. Allora si ha $64k^2 = 8(8h + 4)^2 = 8(64h^2 + 64h + 16)$ e dividendo per 64 si ha $k^2 = 8(h^2 + h) + 2$ assurdo poiché (vedi Tabella 1.1) nessun quadrato è congruo a 2 modulo 8.
- Caso $[n] = [4]$: ossia $n = 8k + 4$. Si giunge ancora all'assurdo di un quadrato in classe $[2]$: infatti $(8k + 4)^2 = 64k^2 + 64k + 16 = 8m^2$ implica che $8(k^2 + k) + 2 = m^2$.

Ripetiamo l'esercizio per $a = 12$. Supponiamo che esistano interi positivi n, m primi tra loro tali che $n^2 = 12 \cdot m^2$. Ora da $n^2 = 12 \cdot m^2$ segue che in $\mathbb{Z}_{12}$ si ha $[n^2] = [n]^2 = [0]$, e quindi (vedi Tabella 1.1) possiamo avere due casi: $[n] = [0]$, oppure $[n] = [6]$.

- Caso $[n] = [0]$.
 - Se $[m^2] \neq [0]$: si procede con la solita prova; $12 \cdot m^2$ è divisibile per 12 ma non per 12^2 , mentre n^2 è divisibile per 12^2 .
 - Se $[m^2] = [0]$: sempre dalla Tabella 1.1 ci sono due possibilità:
 - * $[m] = [0]$: si esclude perché n, m sono primi fra loro
 - * $[m] = [6]$: ossia $n = 12k$, $m = 12h + 6$. Allora si ha $144k^2 = 12(12h+6)^2 = 12(144h^2 + 144h + 36)$ e dividendo per 144 si ha $k^2 = 12(h^2 + h) + 3$ assurdo poiché (vedi Tabella 1.1) nessun quadrato è congruo a 3 modulo 12.
- Caso $[n] = [6]$: ossia $n = 12k + 6$. Si giunge ancora all'assurdo di un quadrato in classe $[3]$: infatti $(12k+6)^2 = 144k^2 + 144k + 36 = 12m^2$ implica che $12(k^2 + k) + 3 = m^2$.

Sempre per esercizio, utilizzando ancora il pacchetto R,⁵ verifichiamo che se $a = 1001$ si ha $x^2 \equiv 0 \pmod{1001}$ solo se $x = 1001 \cdot h$ con h intero. Al comando `length(which(mod((0:1000)^2, 1001)==0))`, che chiede *quanti* dei quadrati dei numeri fra 0 e 1000 sono congrui a 0 modulo 1001, R risponde “1”, ossia il solo 0. Quindi, per il Teorema 1.2.2, l'intero $a = 1001$ non ha radice quadrata razionale.

Sul metodo effettivamente usato da Teodoro esistono solo delle congetture: in [41] viene suggerito che uno dei metodi possibili fosse proprio quello qui presentato nel Teorema 1.2.2.

1.2.3 La radice quadrata degli interi non quadrati

Si ha il seguente risultato generale:

Teorema 1.2.3 *Se un intero positivo a non è un quadrato, allora non esistono interi positivi n, m tali che $n^2 = a \cdot m^2$.*

⁵ Stiamo facendo quella che in inglese si chiama una *computer assisted proof*, ma in versione *light*, in quanto tutti i conti necessari possono essere eseguiti a carta e matita senza alcun particolare problema.

DIMOSTRAZIONE. La dimostrazione più nota fa uso del teorema fondamentale dell'aritmetica (cf. ad es. [18], p. 65).⁶ Supponiamo esistano interi positivi n, m tali che $n^2 = a \cdot m^2$; si può supporre che tali n, m siano primi fra loro. Nessun fattore primo di n può trovarsi fra quelli di m , quindi devono tutti trovarsi fra quelli di a . Per l'unicità, tali fattori primi di a hanno tutti esponente pari ≥ 2 , ossia a è un quadrato. $\square$

Ma se si esce dagli interi, e si lavora in $\mathbb{Q}$ utilizzando la funzione *floor*, o *parte intera*

$$r \in \mathbb{Q} \mapsto [r] = \max\{k \in \mathbb{N} \mid k \leq r\}$$

esistono varie altre dimostrazioni fra le quali una delle più brevi è la seguente, di Richard Beigel [6]:

DIMOSTRAZIONE. Sia per assurdo r un razionale positivo, non intero, tale che $r^2 = a$. Per essere r razionale, esistono interi $n, m > 0$ tali che $r = n/m$ e quindi mr è un intero. Esiste allora il *minimo* intero $q > 0$ tale che qr sia intero, ed $s = q(r - [r]) = qr - q[r]$ è un intero. Per essere r non intero, si ha $r - [r] > 0$ e pertanto $s > 0$. Per definizione, la “mantissa” $r - [r]$ è < 1 , da cui $s < q$. Abbiamo ora che $sr = (qr - q[r])r = qr^2 - qr[r] = qa - qr[r]$ è un intero, contraddicendo il fatto che come q è stato scelto il minimo intero positivo tale che qr sia intero. $\square$

1.3 L'estremo superiore

Tuttavia tutto questo non deve far pensare che $\mathbb{Q}$ “non basta” perché in $\mathbb{Q}$ non si possono fare le radici quadrate di tutti gli interi positivi. L'esigenza delle radici è di tipo algebrico, mentre le esigenze dell'analisi sono maggiormente topologiche (metriche): $\mathbb{Q}$ non basta all'analisi perché ad esempio in $\mathbb{Q}$ non si riescono a “validare” teoremi fondamentali come l'esistenza di massimo e minimo per funzioni continue su un intervallo compatto (Teorema di Weierstrass). Ad esempio la funzione polinomiale

$$f : x \in [-2, 2] \cap \mathbb{Q} \mapsto -2x + \frac{x^3}{3} \in \mathbb{Q}$$

⁶ Ossia l'esistenza ed unicità della scomposizione in fattori primi di un qualunque intero $a \neq 0, \pm 1$. Si noti che l'unicità non è affatto intuitiva: essa ci dice per esempio, senza fare i conti, che $2 \cdot 7^4 \cdot 13 \neq 2 \cdot 3 \cdot 5^4 \cdot 17$ (questi due interi sono 62426 e 63750).

che opera fra numeri razionali, pur se definita su un intervallo (razionale) chiuso e limitato, non ha massimo né minimo:⁷ vedi Figura 1.1 (si noti il contrasto fra l'intuizione proposta dal disegno ed il testo della didascalia). Il problema vero di $\mathbb{Q}$ è che in $\mathbb{Q}$ non vale il teorema “Ogni sottoinsieme non-vuoto superiormente limitato ha estremo superiore”.⁸

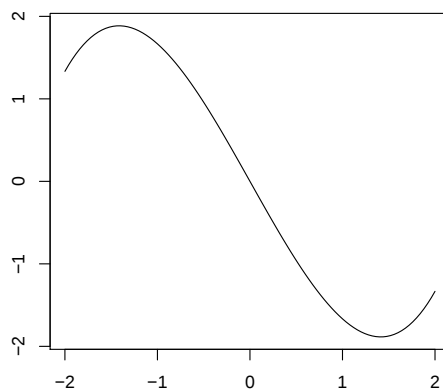


Figura 1.1: Grafico della funzione razionale di variabile razionale definita da $f(x) = -2x + \frac{x^3}{3}$ su $-2 \leq x \leq 2$. La funzione non ha massimo né minimo.

Per comprendere meglio l'esempio che vogliamo fare, conviene immaginare di aver già introdotto sia i numeri reali che la loro rappresentazione decimale (cf. Capitolo 6): rivisitiamo allora il teorema di esistenza dell'estremo superiore in $\mathbb{R}$ dal punto di vista della rappresentazione decimale.

⁷ Diamone una dimostrazione senza uscire da $\mathbb{Q}$. Assumiamo $0 < x < y$ (la funzione è dispari), si ha $f(x) - f(y) = (1/3)(x - y)(-6 + x^2 + xy + y^2)$. Se $y^2 < 2$ e $x^2 < 2$ si ha $2xy \leq x^2 + y^2 < 4$, per cui $f(x) > f(y)$. Se $x^2 > 2$ e $y^2 > 2$ si ha $xy > 2$ (se fosse $y \leq 2/x$ sarebbe $y^2 \leq 4/x^2 < 2$), quindi $f(x) < f(y)$. Dunque f è strettamente decrescente sui razionali $0 < r < 2$ tali che $r^2 < 2$, ed è strettamente crescente sui razionali $0 < r < 2$ tali che $r^2 > 2$; non ce ne sono altri di razionali $0 < r < 2$, per cui sui razionali $0 < r < 2$ la funzione non ha minimo.

⁸ Come vedremo in seguito, questo si riassume dicendo che il *corpo ordinato* $\mathbb{Q}$ non è completo.

Sia X un insieme non vuoto di numeri reali *positivi*, superiormente limitato, ad esempio da 1. Scriviamo tutti i numeri di X in base 10, usando nei casi ambigui la rappresentazione illimitata 9-periodica e non la decimale limitata.⁹ I numeri di X sono quindi del tipo $x = 0.a_1a_2a_3 \dots$

- Sia D_1 l'insieme delle cifre dei *decimi* degli elementi di X , e sia ad esempio $7 = \max(D_1)$ (si noti che le cifre dei decimi costituiscono un insieme finito, che ha al massimo 10 elementi); definiamo (provvisoriamente) il numero $\xi = 0.7$.
- consideriamo il sottoinsieme $X_1 \subseteq X$ costituito dai numeri di X con cifra dei decimi uguale a 7, ossia del tipo $x = 0.7\dots$. Sia D_2 l'insieme delle cifre dei *centesimi* degli elementi di X_1 , e sia ad esempio $4 = \max(D_2)$ (si noti che le cifre dei centesimi costituiscono un insieme finito, che ha al massimo 10 elementi); ridefiniamo allora $\xi = 0.74$.
- consideriamo il sottoinsieme $X_2 \subseteq X_1$ costituito dai numeri di X il cui sviluppo decimale inizia con 74, ossia del tipo $x = 0.74\dots$. Sia D_3 l'insieme delle cifre dei *millesimi* degli elementi di X_2 , e sia per esempio $4 = \max(D_3)$; ridefiniamo allora $\xi = 0.744$; e così via.

Induttivamente, in notazione decimale, rimane definito il numero $\xi = 0.744\dots$ e risulta

$$\xi = 0.744\dots = \sup(X)$$

Infatti tale ξ ha le seguenti due proprietà:

1. Il numero ξ è maggiore o uguale di ogni elemento $x \in X$. Questo è vero per costruzione: ogni $x \in X$ ha cifra dei decimi ≤ 7 ; quelli che iniziano con 0, 1, 2, 3, 4, 5, 6 sono ovviamente $< \xi$; quelli che iniziano con 7 hanno cifra dei centesimi ≤ 4 ; ecc.
2. Per ogni $y < \xi$, esiste un elemento $x \in X$ tale che $x \geq y$. Se $y < \xi$, esiste una prima posizione dopo la virgola dove la cifra di y è < della corrispondente cifra di ξ (altrimenti i due numeri sono uguali). Per esempio potrebbe essere $y = 0.743\dots$. Allora y è superato da tutti i numeri $x \in X$ che iniziano con 744 (e ce ne sono).

⁹ Rappresentiamo cioè, ad esempio, il numero $3/2$ non nella forma 1.5, ma nella forma 1.499999...

Teorema 1.3.1 *In $\mathbb{Q}$ non è vero che ogni insieme non-vuoto superiormente limitato ha estremo superiore.*

DIMOSTRAZIONE. Definiamo un insieme di razionali scritti in base 10, come segue:¹⁰

$$0, 0.1, 0.101, 0.101001, 0.1010010001, 0.101001000100001, \dots$$

L'insieme A di tali numeri è superiormente limitato, ad esempio da

$$0.\bar{1} = 0.1111111\dots = 1/9$$

Se A avesse un estremo superiore r in $\mathbb{Q}$, tale r non potrebbe essere altro che

$$r = 0.101001000100001000001000000100000001\dots$$

Ma r è irrazionale, in quanto non ha sviluppo decimale periodico, e pertanto A è un sottoinsieme di $\mathbb{Q}$, non vuoto, superiormente limitato, ma non ha estremo superiore in $\mathbb{Q}$. $\square$

Si noti che il Teorema 1.3.1 mostra anche che in $\mathbb{Q}$ non vale il teorema sul limite delle funzioni monotone.

Chiudiamo il Capitolo ricordando una proprietà di $\mathbb{Q}$ che verrà richiamata in seguito:

Proposizione 1.3.2 *$\mathbb{Q}$ non ha automorfismi diversi dall'identità.*

DIMOSTRAZIONE. Se $f : \mathbb{Q} \rightarrow \mathbb{Q}$ è un automorfismo,¹¹ allora $f(0) = 0$ e $f(1) = 1$. Ne segue $f(n) = n$ sugli interi, e infine $f(n/m) = f(n)/f(m) = n/m$. $\square$

¹⁰ Zero virgola un 'uno' e uno 'zero', poi un 'uno' e due 'zeri', poi un 'uno' e tre 'zeri', ecc.

¹¹ Si dice che $f : \mathbb{Q} \rightarrow \mathbb{Q}$ è un *automorfismo* se f è un'applicazione biettiva tale che per ogni $x, y \in \mathbb{Q}$, si ha $f(x + y) = f(x) + f(y)$, ed $f(xy) = f(x)f(y)$ (cf. [26] p. 258).

Capitolo 2

Definizione assiomatica di $\mathbb{R}$

On n'a pas, je crois, tiré de l'axiomatique hilbertienne la vraie leçon qui s'en dégage; c'est celle-ci: on n'accède à la rigueur absolue qu'en éliminant la signification; l'absolue rigueur n'est possible que dans et par l'insignifiance. Mais s'il faut choisir entre rigueur et sens, je choisirai sans hésitation le sens.^a

RENÉ THOM,

Mathématiques modernes, mathématiques de toujours

^a Non si è ricavata, credo, dall'assiomatica di Hilbert la vera lezione in essa contenuta: si accede al rigore assoluto solo eliminando il significato; il rigore assoluto è possibile soltanto e per mezzo di tale abbandono di significato. Ma se si deve scegliere tra rigore e significato, scelgo quest'ultimo senza esitare.

2.1 Introduzione

In questo capitolo presentiamo l'introduzione assiomatica di $\mathbb{R}$ come campo ordinato *completo*, cioè come un campo ordinato in cui ogni sottoinsieme non-vuoto e superiormente limitato ammette l'estremo superiore. L'“unicità” di una tale struttura viene dimostrata, in modo classico, applicando un teorema dovuto a David Hilbert (Teorema 2.5.3); cioè viene provato che due campi ordinati completi sono ordinatamente isomorfi, e

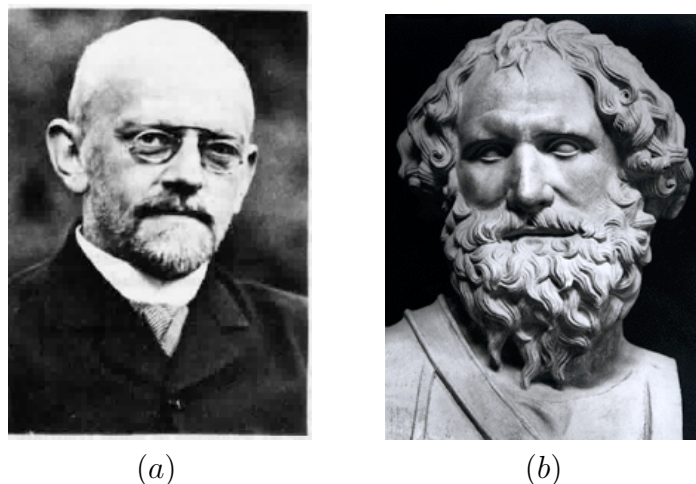


Figura 2.1: (a) D. Hilbert (1862–1943); (b) Archimede (ca. 287–212 a.C.).

pertanto esiste “un unico” sistema dei numeri reali (Teorema 2.5.7).¹ Si prosegue presentando la definizione storica di $\mathbb{R}$ data da Hilbert, basata sull’*Axiom der Vollständigkeit*, e se ne dimostra l’equivalenza con quella assiomatica attuale (Teorema 2.6.12). Questa parte, per la quale gli autori non hanno trovato riferimenti in letteratura, mostra in particolare che la Proprietà di Archimede (ossia l’illimitatezza dei numeri naturali nei reali) svolge un ruolo essenziale nella definizione di $\mathbb{R}$. Alla fine del capitolo, per le implicazioni didattiche, discutiamo le definizioni assiomatiche di $\mathbb{R}$ basate sulle *proprietà di separazione*: la Proprietà di Cantor (esistenza di elemento separatore fra classi contigue) e la Proprietà di Dedekind (esistenza di elemento separatore fra classi separate), con le sue diverse varianti presenti nei libri di testo.

¹ Questa “unicità” del sistema dei numeri reali venne enunciata da David Hilbert nel 1900 nel suo articolo [29]. Secondo J. Kennedy [32], Hilbert avrebbe dato indicazioni per la dimostrazione in alcune sue lezioni della fine del decennio precedente (non pubblicate); successivamente, verso la fine degli anni ’20 del Novecento, l’“unicità” del sistema dei numeri reali divenne un risultato popolarmente noto (“folklore”, [32]). Rimandiamo a [37] per approfondimenti su questo tema.

2.2 Campi ordinati

Ricordiamo che (cf. [26], p. 136; [18], pp. 98, 96, 25):

1. Una relazione $\mathcal{R}$ in un insieme non vuoto X si dice *relazione di ordine totale* se:
 - a) $x\mathcal{R}x$ per ogni $x \in X$ (proprietà riflessiva);
 - b) se $x\mathcal{R}y$ e $y\mathcal{R}x$ allora $x = y$ (proprietà antisimmetrica);
 - c) se $x\mathcal{R}y$ e $y\mathcal{R}z$ allora $x\mathcal{R}z$ (proprietà transitiva);
 - d) per ogni $x, y \in X$, $x \neq y$, si ha $x\mathcal{R}y$ oppure $y\mathcal{R}x$.
2. Un *corpo* è una terna $(\mathbb{K}, +, \cdot)$ dove $\mathbb{K}$ è un insieme non vuoto, “+” e “ $\cdot$ ” sono operazioni binarie su $\mathbb{K}$ tali che:
 - a) $(\mathbb{K}, +)$ è un gruppo commutativo con elemento neutro $0_{\mathbb{K}}$;
 - b) $(\mathbb{K} \setminus \{0_{\mathbb{K}}\}, \cdot)$ è un gruppo con elemento neutro $1_{\mathbb{K}}$, con $1_{\mathbb{K}} \neq 0_{\mathbb{K}}$;
 - c) per ogni $a, b, c \in \mathbb{K}$ vale la legge distributiva $a \cdot (b + c) = a \cdot b + a \cdot c$ e $(a + b) \cdot c = a \cdot c + b \cdot c$.
3. Un *corpo commutativo* o *campo* è un corpo in cui vale la legge commutativa dell’operazione “ $\cdot$ ”.
4. Se a è un elemento di un campo $\mathbb{K}$ si ha $(-1_{\mathbb{K}}) \cdot a = -a$; $(-a) \cdot b = a \cdot (-b) = -(a \cdot b)$; $(-a)^2 = a^2$.

Nel seguito, quando non vi sia ambiguità, l’elemento neutro rispetto alla “+” e l’elemento neutro rispetto alla “ $\cdot$ ” di un corpo $(\mathbb{K}, +, \cdot)$ verranno indicati semplicemente con 0 e 1, e chiamati, rispettivamente, “zero” ed “uno”.

Definizione 2.2.1 *Un campo ordinato è un campo $(\mathbb{K}, +, \cdot)$ munito di una relazione d’ordine totale $\leq$ tale che, per ogni $a, b, c \in \mathbb{K}$ sia:*

$$a \leq b \Rightarrow a + c \leq b + c \quad (2.1)$$

$$0 \leq c, a \leq b \Rightarrow c \cdot a \leq c \cdot b \quad (2.2)$$

Osservazione 2.2.2 In un campo ordinato si pone

$$a < b \iff (a \neq b \text{ e } a \leq b)$$

Se $\mathbb{K}$ è un campo ordinato, l'elemento a si dice *positivo* se $0 < a$, si dice *negativo* se $a < 0$. Di seguito, scriveremo anche $a \geq b$ quando $b \leq a$, e $a > b$ quando $b < a$.

L'insieme $\mathbb{K}^+ = \{a \in \mathbb{K} \mid a > 0\}$ degli elementi positivi viene detto *cono positivo*. Dalla (2.2) con $a = 0$ e $b = c > 0$ segue che se $b > 0$ anche $b^2 > 0$; inoltre, avendosi $(-b)^2 = b^2$, si ha che il quadrato di un qualunque elemento non nullo è positivo. In particolare $1 = 1^2 > 0$. Si osservi che:

(P1) Se $a, b \in \mathbb{K}^+$, allora $a + b, ab \in \mathbb{K}^+$.

(P2) Se $b \neq 0$, vale una ed una sola delle relazioni: $b \in \mathbb{K}^+$ oppure $-b \in \mathbb{K}^+$.

Per provare la (P1), siano $a, b > 0$; da $a > 0$ segue per la (2.1) che $a + b > 0 + b = b > 0$; inoltre, supponendo $a \geq b$, per la (2.2) si ha $ab \geq bb = b^2 > 0$. Per provare la (P2) supponiamo $b \neq 0$ e $b \notin \mathbb{K}^+$, ossia $b < 0$; dalla (2.1) si ha $b + (-b) < 0 + (-b)$, $0 < -b$ e dunque $-b \in \mathbb{K}^+$. $\square$

L'osservazione precedente fornisce un modo pratico per “ordinare” un campo nel caso in cui questo sia possibile: basta definire quali sono gli elementi “positivi”. Più precisamente:

Osservazione 2.2.3 Se in un campo $\mathbb{K}$ esiste un sottoinsieme non vuoto $\mathbb{K}^+$ che verifichi le precedenti proprietà (P1)–(P2) e tale che $0 \notin \mathbb{K}^+$, allora ponendo

$$\begin{cases} a < b & \iff b - a \in \mathbb{K}^+ \\ a \leq b & \iff ((a < b) \text{ oppure } (a = b)) \end{cases}$$

si ottiene un ordine totale che fa di $\mathbb{K}$ un campo ordinato. Le verifiche sono banali.²

² Non tutti i campi possono essere “ordinati” nel senso della Definizione 2.2.1. L'esempio classico è il campo $\mathbb{C}$ dei numeri complessi, nel quale non può essere definita alcuna relazione d'ordine totale che ne faccia un campo ordinato. Infatti, se $\mathbb{C}$ fosse un campo ordinato, dato che in ogni campo ordinato il quadrato di un qualunque elemento non nullo è positivo, si avrebbe che $-1 = i^2$ e $1 = 1^2$ sono entrambi positivi, contro la (P2).

Esempio 2.2.4 Il campo $\mathbb{Q}$ dei numeri razionali è il *campo delle frazioni* o *campo quoziente* sul dominio di integrità $\mathbb{Z}$ (cf. [18] pp. 98, 236). Sin dalla scuola media inferiore si è usato in $\mathbb{Q}$ l'ordinamento canonico, o naturale, per il quale, ad esempio

$$\frac{2}{7} < \frac{3}{4}$$

Mostriamo che a tale ordinamento non ci sono alternative: ossia esiste *un'unica* relazione d'ordine totale su $\mathbb{Q}$ che estende quella di $\mathbb{Z}$ e lo rende un campo ordinato.

Infatti, poiché *in ogni* campo ordinato $b \neq 0$ implica $b^2 > 0$, se un quoziente a/b è positivo anche il prodotto $(a/b)b^2 = ab$ è positivo e viceversa. Quindi in $\mathbb{Q}$ si ha:

$$a/b > 0 \iff ab > 0$$

con $a, b \in \mathbb{Z}$, $b \neq 0$. In particolare, se $a, b, c, d \in \mathbb{Z}$, $b > 0$, $d > 0$ risulta:

$$a/b < c/d \iff c/d - a/b = (bc - ad)/(db) > 0 \iff ad < bc \quad \square$$

In generale il campo $\mathbb{F}$ delle frazioni su un dominio $\mathbb{D}$ di integrità *ordinato*, ossia munito di una relazione d'ordine totale $\leq$ che verifica le (2.1)–(2.2), può essere ordinato con il metodo dell'Osservazione 2.2.3 ponendo $a/b > 0$ in $\mathbb{F}$ se e solo se $ab > 0$ in $\mathbb{D}$. Ritorneremo fra poco su questo punto.

Osservazione 2.2.5 Abbiamo già visto che in un campo ordinato si ha $1 > 0$, quindi, per la (2.1) si ha $1 + 1 > 1$, ed in generale per ogni $n \in \mathbb{N}$, $n > 0$, si ha

$$n1 = \underbrace{1 + 1 + 1 + \dots + 1}_{n \text{ addendi}} > 0$$

Definizione 2.2.6 Si dice che il campo $\mathbb{K}$ ha caratteristica zero se per ogni $n \in \mathbb{N}$, $n > 0$, si ha $n1 \neq 0$.

Segue allora immediatamente dalla Osservazione 2.2.5 che:

Proposizione 2.2.7 Ogni campo ordinato ha caratteristica zero.

Quindi un campo ordinato è necessariamente infinito e conseguentemente i campi finiti $\mathbb{Z}_p$ (cf. [26] p. 132, 392) non possono essere campi ordinati.

Il campo $\mathbb{Q}$ dei numeri razionali, essendo ordinato, ha quindi caratteristica zero. In effetti, $\mathbb{Q}$ è il “più piccolo” campo con caratteristica zero. Questo è il senso della seguente:

Proposizione 2.2.8 *Sia $\mathbb{K}$ un campo di caratteristica zero. Allora esiste un unico omomorfismo di campi³ $\Psi : \mathbb{Q} \rightarrow \mathbb{K}$, e tale omomorfismo risulta iniettivo.*

DIMOSTRAZIONE. La definizione di Ψ è ovvia:

$$\Psi : \frac{p}{q} \in \mathbb{Q} \mapsto \frac{p \cdot 1_{\mathbb{K}}}{q \cdot 1_{\mathbb{K}}} \in \mathbb{K}$$

□

Nel seguito, dato un campo ordinato $\mathbb{K}$ identificheremo il suo sottocampo $\Psi(\mathbb{Q})$ isomorfo a $\mathbb{Q}$ con $\mathbb{Q}$ stesso. Il sottocampo $\Psi(\mathbb{Q})$ è detto il *sottocampo razionale* di $\mathbb{K}$, ed i suoi elementi sono detti gli *elementi razionali* di $\mathbb{K}$. Osserviamo ancora che se $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ è un omomorfismo fra due campi ordinati $\mathbb{K}_1$ e $\mathbb{K}_2$, si ha necessariamente

$$\Phi(p/q) = p/q \quad (2.3)$$

(dove il primo p/q scritto è un elemento di $\mathbb{K}_1$ ed il secondo p/q scritto è un elemento di $\mathbb{K}_2$), ossia un omomorfismo fra due campi ordinati trasforma un razionale $p/q \in \mathbb{K}_1$ in $p/q \in \mathbb{K}_2$; si usa dire per brevità che Φ è la identità su $\mathbb{Q}$. Ad esempio:

$$\Phi(\underbrace{3/2}_{\in \mathbb{K}_1}) = \Phi\left(\frac{1_{\mathbb{K}_1} + 1_{\mathbb{K}_1} + 1_{\mathbb{K}_1}}{1_{\mathbb{K}_1} + 1_{\mathbb{K}_1}}\right) = \frac{1_{\mathbb{K}_2} + 1_{\mathbb{K}_2} + 1_{\mathbb{K}_2}}{1_{\mathbb{K}_2} + 1_{\mathbb{K}_2}} = \underbrace{3/2}_{\in \mathbb{K}_2}$$

³ Dati i campi $(\mathbb{K}, +, \cdot)$ e $(\mathbb{K}', +, \cdot)$, un'applicazione $f : \mathbb{K} \rightarrow \mathbb{K}'$ è un *omomorfismo* se per ogni $x, y \in \mathbb{K}$ risulta $f(x+y) = f(x) + f(y)$ ed $f(x \cdot y) = f(x) \cdot f(y)$. Se f è un omomorfismo fra due campi si ha necessariamente che $f(0) = 0$, che $f(1) = 1$ e che $f(x) = f(y)$ implica $x = y$, cioè un omomorfismo fra campi è sempre *iniettivo*. Infatti, sia f un omomorfismo fra campi, e supponiamo per assurdo che risulti $f(x) = f(y)$ con $x \neq y$: allora, posto $u = x - y$, si ha $f(u) = 0$ con $u \neq 0$; pertanto, esiste l'inverso u^{-1} e si ha, da un lato, $f(u \cdot u^{-1}) = f(1) = 1$, dall'altro, $f(u \cdot u^{-1}) = f(u) \cdot f(u^{-1}) = 0$, ciò che contraddice l'assunto $1 \neq 0$. Un omomorfismo f fra campi che sia biiettivo è detto un *isomorfismo*.

2.3 Campi ordinati archimedei

Definizione 2.3.1 *Un campo ordinato $\mathbb{K}$ si dice archimedeo se vale la cosiddetta Proprietà di Archimede: per ogni $x \in \mathbb{K}$, $x > 0$, esiste $n \in \mathbb{N}$ tale che $n \cdot 1 > x$.*

Dalla definizione segue che, se $\mathbb{K}$ è un campo ordinato archimedeo, allora per ogni $a \in \mathbb{K}$, $a > 0$ e per ogni $b \in \mathbb{K}$, $b > 0$, esiste $n \in \mathbb{N}$ tale che $n \cdot b > a$; basta infatti considerare $n \in \mathbb{N}$ tale che $n \cdot 1 > a \cdot b^{-1}$ e si ottiene la disuguaglianza cercata.⁴

Esempio 2.3.2 *Esistono campi ordinati archimedei.* Il campo $\mathbb{Q}$ dei numeri razionali è evidentemente archimedeo, in quanto ogni numero razionale può essere rappresentato come n/m con n, m interi ed $m > 0$ e si ha $|n| \cdot 1 > n/m$.

Produrre invece un esempio di campo ordinato non-archimedeo è semplice, ma non altrettanto banale.

Esempio 2.3.3 *Esistono campi ordinati non-archimedei.* Ricordiamo (cf. [18] p. 253) che l'anello $A[x]$ dei polinomi in x a coefficienti in un dominio di integrità A è anch'esso un dominio di integrità. Per esempio, l'anello $\mathbb{Z}[x]$ dei polinomi a coefficienti interi è un dominio d'integrità. Ora i polinomi a coefficienti interi possono essere ordinati indicando come positivi quei polinomi con coefficiente direttore⁵ positivo. Il campo $\mathbb{F}$ delle frazioni su $\mathbb{Z}[x]$, ossia il campo i cui elementi sono le funzioni razionali fratte a coefficienti interi $p(x)/q(x)$, può essere ordinato canonicamente ponendo

$$p(x)/q(x) > 0 \iff p(x)q(x) > 0$$

Praticamente, se standardizziamo la rappresentazione di queste frazioni assumendo positivo il coefficiente direttore del denominatore, abbiamo

⁴ La Proprietà di Archimede ricorda il proverbio “a quattrino a quattrino si supera lo zecchino”. Il riferimento esplicito ad Archimede venne introdotto nel 1883 dal matematico austriaco Otto Stolz [54], in quanto Archimede, nella sua opera *Arenario* (*Ψαμμίτης*, in lat.: *Arenarius*), mostra che esiste un numero intero n così grande che n granelli di sabbia riempirebbero la sfera dell'universo allora conosciuto, ossia che, parafrasando il proverbio, “a granello a granello si supera l'universo”.

⁵ Il coefficiente direttore di un polinomio è il coefficiente (non nullo) del termine di grado massimo.

semplicemente che sono positive quelle frazioni che, così standardizzate, hanno positivo il coefficiente direttore del numeratore. Inoltre si ha $p(x)/q(x) > r(x)/s(x)$ se e solo se $p(x)/q(x) - r(x)/s(x) > 0$. Ad esempio,

$$\frac{3x^2 + 1}{4x^2 + 5} > \frac{2x^2 + x + 1}{5x^2 + 1}$$

in quanto

$$\frac{3x^2 + 1}{4x^2 + 5} - \frac{2x^2 + x + 1}{5x^2 + 1} = \frac{-4 - 5x - 6x^2 - 4x^3 + 7x^4}{5 + 29x^2 + 20x^4} > 0$$

Il campo $\mathbb{F}$ non è archimedeo, perché nessun multiplo n della frazione “costante” uguale ad $1/1$ può superare la frazione $x/1$. Infatti, per ogni n intero, $n \geq 0$, si ha che

$$\frac{x}{1} > \frac{n}{1}$$

perché risulta

$$\frac{x}{1} - \frac{n}{1} = \frac{x - n}{1} > 0$$

In generale le frazioni $p(x)/q(x)$ con $\deg(p) > \deg(q)$ sono dette *infinite*, nel senso che sono maggiori di qualunque numero naturale n . Invece le frazioni $p(x)/q(x) > 0$ con $\deg(p) < \deg(q)$ sono dette *infinitesime*, nel senso che sono minori di $1/n$, per qualunque numero naturale $n > 0$. Ad esempio:

$$\begin{aligned} \text{(a)} \quad \frac{x^2 + 1}{2x + 5} > n \quad \text{perché} \quad \frac{x^2 + 1}{2x + 5} - n &= \frac{1 - 5n - 2nx + x^2}{5 + 2x} > 0 \\ \text{(b)} \quad \frac{x + 1}{2x^2 + 5} < \frac{1}{n} \quad \text{perché} \quad \frac{1}{n} - \frac{x + 1}{2x^2 + 5} &= \frac{5 - n - nx + 2x^2}{5n + 2nx^2} > 0 \end{aligned}$$

La Proprietà di Archimede – se vale in un campo ordinato – esclude che esistano elementi infiniti ed esclude che esistano elementi infinitesimi.

Proposizione 2.3.4 *Sia $\mathbb{K}$ un campo ordinato archimedeo. In ogni intervallo aperto $]a, b[$ con $a, b \in \mathbb{K}$, $a < b$, esiste un elemento razionale.*

DIMOSTRAZIONE. Supponiamo inizialmente $0 < a < b$. Dato che $b - a > 0$ la Proprietà di Archimede implica che esiste $n \in \mathbb{N}$, $n > 0$, tale che $n \cdot (b - a) > 1$ ossia $b - a > 1/n$. L'insieme

$$A = \{k \in \mathbb{N} \mid k > 0, k/n > a\}$$

è non vuoto: questo segue dalla Proprietà di Archimede applicata agli elementi positivi $1/n$ e a . Sia $m = \min A$ (ogni sottoinsieme non vuoto di $\mathbb{N}$ ha minimo). Dato che $m \in A$, si ha $m/n > a$. Poiché m è il minimo di A , $m-1$ non appartiene ad A , pertanto $(m-1)/n \leq a$, da cui, usando anche il fatto che $1/n < b-a$,

$$\frac{m}{n} = \frac{m-1}{n} + \frac{1}{n} \leq a + \frac{1}{n} < a + (b-a) = b$$

Quindi $m/n \in]a, b[$, come richiesto. Possiamo ricondurre al caso $a > 0$ il caso di un intervallo qualsiasi $]a, b[$, infatti, se $a \leq 0$, basta considerare $r \in \mathbb{N}$ tale che $a+r > 0$. Come dimostrato, esiste un razionale $m/n \in]a+r, b+r[$, e pertanto $m/n - r$ è un elemento razionale appartenente all'intervallo $]a, b[$. $\square$

Definizione 2.3.5 *Sia A un insieme totalmente ordinato. Un sottoinsieme B di A si dice denso in A se ha intersezione non vuota con ogni intervallo aperto di A .*

Osservazione 2.3.6 La Proposizione 2.3.4 afferma che $\mathbb{Q}$ è denso in ogni campo ordinato archimedeo. Iterando questa proposizione, si verifica facilmente che *ogni intervallo aperto di un campo ordinato archimedeo contiene infiniti elementi razionali.*

Osservazione 2.3.7 In seguito avremo bisogno di un criterio che ci consenta di stabilire se un elemento a di un campo ordinato archimedeo $\mathbb{K}$ è uguale a 0. Un criterio utile è il seguente:

Dato un campo ordinato archimedeo $\mathbb{K}$ e dato $a \in \mathbb{K}$, se per ogni $n \in \mathbb{N}$, $n > 0$, risulta $|a| \leq 1/n$, allora $a = 0$.⁶

DIMOSTRAZIONE. Se $a \neq 0$, risulta $|a| > 0$, e, per la Proprietà di Archimede, esiste $n \in \mathbb{N}$ tale che $n \cdot |a| > 1$, ossia $|a| > 1/n$. $\square$

⁶ Abbiamo qui usato il *valore assoluto*, definito al solito modo:

$$|a| = \begin{cases} a & \text{se } a > 0 \\ -a & \text{se } a < 0 \\ 0 & \text{se } a = 0 \end{cases}$$

2.4 Campi ordinati completi

La seguente definizione è fondamentale.⁷

Definizione 2.4.1 *Un campo ordinato $\mathbb{K}$ è completo se per ogni sottoinsieme $A \subseteq \mathbb{K}$ non vuoto e limitato superiormente esiste l'estremo superiore di A .*

Osservazione 2.4.2 Ricordiamo che $\xi = \sup(A)$ se e solo se valgono (simultaneamente) le seguenti due proprietà caratteristiche:

I Per ogni $x \in A$ si ha $x \leq \xi$.

II Per ogni $\epsilon \in \mathbb{K}$, $\epsilon > 0$ esiste $x \in A$ tale che $x > \xi - \epsilon$.

Infatti la I proprietà significa che ξ è un maggiorante di A , e la II proprietà significa che un qualsiasi elemento y di $\mathbb{K}$ che sia minore di ξ (e che quindi può essere scritto come $y = \xi - \epsilon$ con $\epsilon > 0$) non può essere un maggiorante di A , dato che è superato da un elemento $x \in A$. $\square$

Notiamo che, se un campo ordinato $\mathbb{K}$ è *completo*, allora per ogni sottoinsieme $A \subseteq \mathbb{K}$ non vuoto e limitato inferiormente esiste l'estremo inferiore di A : infatti $-A = \{a \in \mathbb{K} \mid -a \in A\}$ è limitato superiormente e si verifica facilmente che

$$\inf A = -\sup(-A)$$

Osservazione 2.4.3 Per evitare i rischi di fraintendimento relativi alla molteplicità di significati dell'aggettivo “completo” già discussi nell'Introduzione, sarebbe opportuno utilizzare nella Definizione 2.4.1 il termine *completo relativamente all'ordinamento* (in inglese: *order-complete*) in luogo di *completo* tout court. Abbiamo scelto di evitare questa pedanteria, ma riteniamo comunque conveniente enfatizzare che, nel presente testo, *campo ordinato completo* significa *campo ordinato completo relativamente all'ordinamento*, ovvero *campo ordinato completo in cui ogni sottoinsieme non vuoto e limitato superiormente ammette l'estremo superiore*. Analogamente, parlando di *completezza* tout court per un campo ordinato, intenderemo sempre *completezza relativamente all'ordinamento*. Si veda anche la precisazione terminologica nella Sezione 2.6.5.

⁷ Useremo qui e nel seguito i concetti di *maggiorante*, *minorante*, *estremo superiore*, *estremo inferiore*, *sottoinsieme limitato inferiormente*, *sottoinsieme limitato superiormente*, ecc.: cf. [22], p. 16–17; [18], p. 26.

Osservazione 2.4.4 È bene osservare che fra gli assiomi che definiscono un campo ordinato completo $\mathbb{K}$, dal punto di vista logico vi è un'importante differenza fra l'assioma di completezza e tutti gli altri. Questi altri, quelli di campo ordinato, hanno i quantificatori $\forall$ ed $\exists$ applicati esclusivamente ad *elementi* di $\mathbb{K}$. La completezza invece richiede il quantificatore universale $\forall$ applicato a *sottoinsiemi* di $\mathbb{K}$. La completezza di $\mathbb{K}$, scritta formalmente, si esprime come segue:

$$\begin{aligned} \forall u \subseteq \mathbb{K} \quad & ((u \neq \emptyset \wedge \exists z \forall x (x \in u \rightarrow x \leq z)) \\ & \rightarrow \exists z (\forall x (x \in u \rightarrow x \leq z) \wedge \forall z' (\forall x (x \in u \rightarrow x \leq z') \rightarrow z \leq z')) \end{aligned}$$

dove u indica un sottoinsieme di $\mathbb{K}$ mentre le altre lettere minuscole indicano elementi di $\mathbb{K}$.

Proposizione 2.4.5 *Ogni campo ordinato completo è archimedeo.*

DIMOSTRAZIONE. Supponiamo per assurdo che esistano $a, b \in \mathbb{K}, a > 0, b > 0$, tali che per ogni $n \in \mathbb{N}, n > 0$ valga $na \leq b$. Allora l'insieme

$$A = \{na \mid n \in \mathbb{N}, n > 0\}$$

è limitato superiormente (b è un maggiorante). Sia $s = \sup A$. Dato che $s - a < s = \sup A$, per la II proprietà caratteristica dell'estremo superiore esiste un elemento $na \in A$, con $n \in \mathbb{N}, n > 0$, tale che $na > s - a$, ossia $(n + 1)a > s$. Ma $(n + 1)a$ è ancora un elemento di A , e l'ultima disuguaglianza contraddice il fatto che s sia un maggiorante di A . $\square$

2.5 Teorema di Hilbert

Giunti a questo punto della teoria, sappiamo che esistono campi ordinati archimedei, ad esempio $\mathbb{Q}$, ma non abbiamo ancora presentato un teorema di esistenza di campi ordinati completi. Conviene come prima cosa dimostrare l'importante Teorema di Hilbert (il Teorema 2.5.3), che mostra una fondamentale proprietà di cui gode un campo ordinato completo rispetto ai campi ordinati archimedei non completi.

2.5.1 Omomorfismi ordinati

Ricordiamo che, come un *campo* va definito assegnando la terna $(\mathbb{K}, +, \cdot)$ costituita dall'insieme sostegno $\mathbb{K}$ e dalle leggi di composizione interna

(nell'ordine: addizione e moltiplicazione), così un *campo ordinato* va formalmente definito assegnando la quaterna $(\mathbb{K}, +, \cdot, \leq)$ costituita dal sostegno, dalle leggi di composizione interna e dalla relazione d'ordine.

Analogamente, come un *omomorfismo fra campi* è un'applicazione $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ che conserva le operazioni interne, così un *omomorfismo fra campi ordinati* è un'applicazione $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ che conserva sia le operazioni interne *sia la relazione d'ordine*, ossia tale che per ogni $a, b \in \mathbb{K}_1$ non solo si ha che $\Phi(a + b) = \Phi(a) + \Phi(b)$ e $\Phi(a \cdot b) = \Phi(a) \cdot \Phi(b)$, ma si ha anche che $a \leq b$ implica $\Phi(a) \leq \Phi(b)$.

Enfatizziamo questo fatto con la seguente:

Definizione 2.5.1 *Se $\mathbb{K}_1$ e $\mathbb{K}_2$ sono campi ordinati, chiameremo omomorfismo ordinato un'applicazione $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ tale che per ogni $a, b \in \mathbb{K}_1$ si ha*

$$\begin{cases} \Phi(a + b) &= \Phi(a) + \Phi(b) \\ \Phi(a \cdot b) &= \Phi(a) \cdot \Phi(b) \\ a \leq b &\Rightarrow \Phi(a) \leq \Phi(b) \end{cases}$$

Un omomorfismo ordinato Φ fra campi ordinati è un'applicazione *strettamente crescente*, cioè da $a < b$ segue che $\Phi(a) < \Phi(b)$, in quanto un omomorfismo fra campi è necessariamente iniettivo. Se un omomorfismo ordinato $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ è biiettivo, si dice che Φ è un *isomorfismo ordinato* e che i campi $\mathbb{K}_1$ e $\mathbb{K}_2$ sono *ordinatamente isomorfi*.

Nel seguito useremo il seguente risultato.

Proposizione 2.5.2 *Sia $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ un'applicazione crescente fra due campi ordinati archimedei. Se per ogni r razionale si ha $\Phi(r) = r$, allora Φ è strettamente crescente.*

DIMOSTRAZIONE. Dati $x < y$ in $\mathbb{K}_1$, si deve provare che $\Phi(x) < \Phi(y)$. Per ipotesi si ha $\Phi(x) \leq \Phi(y)$. Supponiamo per assurdo $\Phi(x) = \Phi(y)$. Per la Proposizione 2.3.4 esistono razionali r ed s , $r \neq s$, tali che $x < r < s < y$. Si ha $\Phi(r) = r$, $\Phi(s) = s$, ed essendo Φ crescente risulta $\Phi(x) \leq r < s \leq \Phi(y)$, da cui $\Phi(x) < \Phi(y)$. $\square$

Teorema 2.5.3 (Teorema di Hilbert) *Sia $\mathbb{K}_1$ un campo ordinato archimedeo, e sia $\mathbb{K}_2$ un campo ordinato completo. Allora esiste un unico omomorfismo ordinato $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$.*

DIMOSTRAZIONE. Vogliamo costruire l'omomorfismo $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$. Ricordiamo che sia $\mathbb{K}_1$ che $\mathbb{K}_2$ contengono una "copia" di $\mathbb{Q}$. Dato $x \in \mathbb{K}_1$, consideriamo il sottoinsieme di $\mathbb{Q}$,

$$A(x) = \{q \in \mathbb{Q} \mid q \leq x\}$$

dove $\leq$ è la relazione d'ordine in $\mathbb{K}_1$. L'insieme $A(x)$ è limitato superiormente in $\mathbb{K}_1$ (x è un maggiorante). Per la Proprietà di Archimede esistono razionali maggiori di x , dunque $A(x)$ è limitato superiormente anche in $\mathbb{Q}$. Consideriamo adesso $A(x)$ come sottoinsieme di $\mathbb{K}_2$. Poiché $A(x)$ è limitato superiormente in $\mathbb{Q}$ e $\mathbb{Q} \subseteq \mathbb{K}_2$, si ha che $A(x)$ è limitato superiormente anche in $\mathbb{K}_2$. Per la completezza di $\mathbb{K}_2$ possiamo definire $\Phi(x) = \sup A(x) \in \mathbb{K}_2$, dove $\sup$ indica l'estremo superiore in $\mathbb{K}_2$. Se $x \in \mathbb{Q}$, l'insieme $A(x)$ ha massimo x , dunque $\Phi(x) = x$. Inoltre, se $x, y \in \mathbb{K}_1$, $x < y$, si ha $A(x) \subseteq A(y)$, da cui

$$\Phi(x) = \sup A(x) \leq \sup A(y) = \Phi(y)$$

e quindi l'applicazione Φ è crescente. Per la Proposizione 2.5.2 la Φ è strettamente crescente.

Rimane da verificare che Φ è un omomorfismo di campi (che sarà necessariamente un omomorfismo ordinato, visto che è strettamente crescente).

Iniziamo con il dimostrare che Φ conserva la somma. Siano $x, y \in \mathbb{K}_1$. Dato $\epsilon \in \mathbb{Q}$, $\epsilon > 0$, esistono $p_1, p_2, q_1, q_2 \in \mathbb{Q}$ tali che

$$x - \epsilon < p_1 < x < p_2 < x + \epsilon, \quad y - \epsilon < q_1 < y < q_2 < y + \epsilon.$$

Quindi

$$p_1 + q_1 < x + y < p_2 + q_2, \quad p_2 - \epsilon < x < p_1 + \epsilon, \quad q_2 - \epsilon < y < q_1 + \epsilon.$$

Dato che Φ è crescente e vale l'identità su $\mathbb{Q}$, troviamo le disuguaglianze

$$p_1 + q_1 \leq \Phi(x + y) \leq p_2 + q_2, \quad p_2 - \epsilon \leq \Phi(x) \leq p_1 + \epsilon, \quad q_2 - \epsilon \leq \Phi(y) \leq q_1 + \epsilon,$$

dove abbiamo sfruttato il fatto che anche ϵ è un razionale. Combinate assieme queste disuguaglianze implicano

$$\Phi(x) + \Phi(y) - 2\epsilon \leq \Phi(x + y) \leq \Phi(x) + \Phi(y) + 2\epsilon$$

e dato che ciò vale per ogni ϵ razionale positivo, concludiamo che $\Phi(x + y) = \Phi(x) + \Phi(y)$.

Verifichiamo ora che $\Phi(xy) = \Phi(x)\Phi(y)$. Poiché $\Phi(x + y) = \Phi(x) + \Phi(y)$, si ha $\Phi(-x) = -\Phi(x)$ e quindi possiamo ridurci al caso $x > 0$, $y > 0$. Siano $\epsilon, p_1, p_2, q_1, q_2$ come sopra, con ϵ così piccolo che $x - \epsilon > 0$ e $y - \epsilon > 0$. Allora

$$p_1 q_1 < xy < p_2 q_2, \quad p_2 - \epsilon < x < p_1 + \epsilon, \quad q_2 - \epsilon < y < q_1 + \epsilon,$$

e per le proprietà di Φ ,

$$p_1 q_1 \leq \Phi(xy) \leq p_2 q_2, \quad p_2 - \epsilon \leq \Phi(x) \leq p_1 + \epsilon, \quad q_2 - \epsilon \leq \Phi(y) \leq q_1 + \epsilon.$$

Dato che $p_2 - \epsilon > x - \epsilon > 0$ e $q_2 - \epsilon > y - \epsilon > 0$, dalle ultime due disuguaglianze segue che

$$(p_2 - \epsilon)(q_2 - \epsilon) \leq \Phi(x)\Phi(y) \leq (p_1 + \epsilon)(q_1 + \epsilon).$$

Si ottiene quindi

$$\Phi(xy) - \Phi(x)\Phi(y) \leq p_2 q_2 - (p_2 - \epsilon)(q_2 - \epsilon) = \epsilon(p_2 + q_2 - \epsilon) < \epsilon(x + y + \epsilon),$$

$$\Phi(xy) - \Phi(x)\Phi(y) \geq p_1 q_1 - (p_1 + \epsilon)(q_1 + \epsilon) = -\epsilon(p_1 + q_1 + \epsilon) < -\epsilon(x + y + \epsilon),$$

da cui

$$|\Phi(xy) - \Phi(x)\Phi(y)| < \epsilon|x + y| + \epsilon^2.$$

Per l'arbitrarietà del razionale positivo ϵ si deduce $\Phi(xy) = \Phi(x)\Phi(y)$.

Per la unicità, supponiamo che $\Phi' : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ sia un omomorfismo ordinato: allora Φ' , essendo un omomorfismo, deve necessariamente trasformare un elemento razionale $q \in \mathbb{Q} \subseteq \mathbb{K}_1$ in $q \in \mathbb{Q} \subseteq \mathbb{K}_2$. Dato che Φ' è strettamente crescente, non può essere $\Phi'(x) < \sup A(x)$. Se fosse $\Phi'(x) > \sup A(x)$, esisterebbe un razionale q tale che $\sup A(x) < q < \Phi'(x)$. In $\mathbb{K}_1$ dovrebbe essere $q < x$ (in quanto $q \geq x$ implicherebbe $\Phi'(q) = q \geq \Phi'(x)$), e quindi $q \in A(x)$, assurdo. Dunque per ogni $x \in \mathbb{K}_1$ si ha $\Phi'(x) = \sup A(x) = \Phi(x)$. $\square$

Osservazione 2.5.4 Notiamo che il Teorema di Hilbert 2.5.3 ha una struttura analoga ad una delle costruzioni classiche dell'analisi matematica, quella della funzione esponenziale di base $a > 1$ ed esponente *reale* x , ossia $f : \mathbb{R} \rightarrow \mathbb{R}$,

$f(x) = a^x$, a partire dalla definizione della potenza $a^{p/q} = \sqrt[q]{a^p}$ con esponente razionale p/q . Nel dominio $\mathbb{R}$ di f si usa il fatto che i razionali sono un insieme denso in $\mathbb{R}$, e si approssima x con i razionali $p/q < x$. Si considera poi l'insieme

$$P = \{a^{p/q} \mid p/q < x\}$$

contenuto nel codominio $\mathbb{R}$ di f , e, sfruttando il fatto che in tale codominio gli insiemi non-vuoti e superiormente limitati hanno estremo superiore, si pone per definizione

$$a^x = \sup P = \sup \{a^{p/q} \mid p/q < x\}$$

L'analogia con il Teorema di Hilbert 2.5.3 è evidente: il Teorema 2.5.3 vuole costruire un omomorfismo ordinato $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$, ed ha proprio come ipotesi che il dominio $\mathbb{K}_1$ sia *archimedeo* (in modo da poter “approssimare” un generico elemento $x \in \mathbb{K}_1$ con elementi razionali) e che il codominio $\mathbb{K}_2$ sia *completo* (in modo da poter “fare” in $\mathbb{K}_2$ l'estremo superiore). Sotto queste ipotesi, Φ viene costruito con la stessa tecnica della costruzione di a^x : per le proprietà algebriche di un omomorfismo, Φ è “inchiodato” sui razionali, e, per la crescita stretta, Φ *deve* essere definito tramite l'estremo superiore come nella dimostrazione del Teorema 2.5.3.

Osservazione 2.5.5 Reinterpretiamo ora il Teorema di Hilbert 2.5.3 nel linguaggio della teoria delle categorie.⁸ Introduciamo la categoria $\mathcal{C}$ che ha per *oggetti* i campi ordinati archimedei $\mathbb{K}_1, \mathbb{K}_2, \dots$, e per *morfismi* gli omomorfismi ordinati $\mathbb{K}_1 \rightarrow \mathbb{K}_2$. Ricordiamo che i morfismi di $\mathcal{C}$, essendo iniettivi, sono necessariamente *strettamente crescenti*. Ricordiamo ancora che un *oggetto finale*⁹ (o *terminale*) in una categoria $\mathcal{C}$ è un oggetto B di $\mathcal{C}$ tale che, per ogni oggetto X di $\mathcal{C}$, esiste uno ed un solo morfismo $X \rightarrow B$. Confrontando questa definizione di oggetto finale con l'enunciato del Teorema di Hilbert 2.5.3, vediamo che quest'ultimo equivale precisamente a dire che *un campo ordinato completo è un oggetto finale nella categoria dei campi ordinati archimedei*. Si conclude quindi che due corpi ordinati completi sono ordinatamente isomorfi in quanto, in una categoria, gli oggetti finali sono necessariamente isomorfi. Il successivo Teorema 2.5.7 mostra questo risultato in dettaglio.

Osservazione 2.5.6 Il Teorema di Hilbert 2.5.3 considera nell'ipotesi un campo ordinato completo. Ci si può chiedere che senso abbia tale ipotesi

⁸ Per un'introduzione di base ai concetti di categoria, oggetti e morfismi cf. [14], p. 93.

⁹ Cf. [14], p. 97.

quando ancora non sappiamo se esistono campi ordinati completi.¹⁰ In effetti la struttura del Teorema 2.5.3 è di tipo implicazione $\mathcal{P} \implies \mathcal{Q}$, ossia $\neg\mathcal{P} \vee \mathcal{Q}$ (non $\mathcal{P}$ oppure $\mathcal{Q}$). Va quindi letta in questo modo: o non esistono campi ordinati completi (cioè $\mathcal{P}$ è falsa e allora il teorema $\mathcal{P} \implies \mathcal{Q}$ è vero in modo “vuoto”), oppure esistono campi ordinati completi, e allora esiste un unico omomorfismo ecc. Una considerazione del tutto analoga vale per il successivo Teorema 2.5.7 di “unicità”.

2.5.2 “Unicità”

Teorema 2.5.7 *Siano $\mathbb{K}_1$ e $\mathbb{K}_2$ campi ordinati completi. Allora esiste un unico isomorfismo ordinato $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$.*

DIMOSTRAZIONE. Per la Proposizione 2.4.5, $\mathbb{K}_1$ e $\mathbb{K}_2$ sono archimedei. Per il Teorema di Hilbert 2.5.3, esistono un omomorfismo ordinato $\Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ ed un omomorfismo ordinato $\Psi : \mathbb{K}_2 \rightarrow \mathbb{K}_1$. Allora le funzioni composte

$$\Psi \circ \Phi : \mathbb{K}_1 \rightarrow \mathbb{K}_1, \quad \Phi \circ \Psi : \mathbb{K}_2 \rightarrow \mathbb{K}_2$$

sono entrambe omomorfismi ordinati. Per la unicità nel Teorema di Hilbert 2.5.3, $\Psi \circ \Phi$ coincide con l’identità su $\mathbb{K}_1$ e $\Phi \circ \Psi$ coincide con l’identità su $\mathbb{K}_2$. Quindi Φ è biettiva e Ψ è la sua inversa. $\square$

Osservazione 2.5.8 Osserviamo che se avessimo a disposizione un *modello* di campo ordinato completo (come ad esempio quello di Méray–Cantor costruito nel Capitolo 3 o quello di Dedekind costruito nel Capitolo 4), potremmo fornire una dimostrazione alternativa del Teorema 2.5.7 provando che un qualunque campo ordinato completo $\mathbb{K}$ è isomorfo al modello dato, tramite un isomorfismo ordinato.

2.5.3 Definizione assiomatica di $\mathbb{R}$

Avendo dimostrato che fra due qualunque campi ordinati completi esiste un isomorfismo ordinato, possiamo dire che esiste “un solo” campo ordinato completo e dare la seguente definizione assiomatica.

¹⁰ Più realisticamente dovremmo dire: ... quando ancora, giunti a questo punto del percorso di questo testo, dobbiamo far finta di non sapere che esistono campi ordinati completi.

Definizione 2.5.9 *Definiamo il campo $\mathbb{R}$ dei numeri reali come l'unico (a meno di isomorfismi ordinati) campo ordinato completo.*

Notiamo che, a questo punto, resta ancora da discutere il problema dell'*esistenza* di un campo ordinato completo, ossia dell'esistenza di $\mathbb{R}$, come già precisato all'inizio di questa Sezione 2.5, ma abbiamo risolto positivamente il problema dell'*unicità* di $\mathbb{R}$ (a meno di isomorfismi ordinati).

2.6 L'Axiom der Vollständigkeit

La definizione assiomatica del campo dei numeri reali compare espressamente nella celeberrima conferenza di Hilbert al II Congresso Internazionale dei Matematici di Parigi del 1900; essa viene richiamata nell'esposizione del II Problema (vedi [27], [28]):

Gli assiomi dell'aritmetica sono essenzialmente nient'altro che le regole ben conosciute delle operazioni, con l'aggiunta dell'assioma di continuità. Recentemente, le ho raccolte tutte assieme e, nel farlo, ho sostituito l'assioma di continuità da due assiomi più semplici, vale a dire, l'assioma ben noto di Archimede e un nuovo assioma, essenzialmente il seguente: i numeri formano un sistema di oggetti che non ammette alcuna ulteriore estensione [*Erweiterung*] in cui continuino a valere tutti gli altri assiomi (assioma di completezza [*Axiom der Vollständigkeit*]). Sono convinto che debba esser possibile trovare una dimostrazione diretta per la compatibilità degli assiomi aritmetici, per mezzo di uno studio attento e di un'opportuna modifica dei metodi noti di ragionamento nella teoria dei numeri irrazionali.¹¹

¹¹ Die Axiome der Arithmetik sind im wesentlichen nichts anderes als die bekannten Rechnungsgesetze mit Hinzunahme des Axiomes der Stetigkeit. Ich habe sie kürzlich zusammengestellt (Jahresbericht der Deutschen Mathematiker-Vereinigung, Bd. 8, 1900, S. 180) und dabei das Axiom der Stetigkeit durch zwei einfachere Axiome ersetzt, nämlich das bekannte Archimedische Axiom und ein neues Axiom des Inhaltes, daß die Zahlen ein System von Dingen bilden, welches bei Aufrechterhaltung der sämtlichen übrigen Axiome keiner Erweiterung mehr fähig ist (Axiom der Vollständigkeit). Ich bin nun überzeugt, daß es gelingen muß einen direkten Beweis für die Widerspruchslösigkeit der arithmetischen Axiome zu finden, wenn man die bekannten Schlußmethoden in der Theorie der Irrationalzahlen im Hinblick auf das bezeichnete

Gli assiomi “storici” di Hilbert per il sistema dei numeri reali, che indicheremo con $\mathbb{R}_H$, sono quindi: un primo gruppo di (sedici) assiomi, che equivalgono, nel linguaggio attuale, a dire che $\mathbb{R}_H$ è un campo ordinato, un 17-esimo assioma che corrisponde alla Proprietà di Archimede, ed un 18-esimo ed ultimo “assioma di completezza” [*Axiom der Vollständigkeit*], che richiede che il sistema $\mathbb{R}_H$ non abbia alcuna “estensione” [*Erweiterung*] in cui continuino a valere tutti gli assiomi precedenti. Intuitivamente, gli assiomi storici di Hilbert per il sistema dei numeri reali $\mathbb{R}_H$ richiedono quindi che questo sia il campo ordinato archimedeo “più grande possibile”.

A proposito dell’*Axiom der Vollständigkeit* sono però importanti due osservazioni. La prima è che Hilbert non utilizza il termine “*Vollständigkeit*”, ossia *completezza*, nel senso tecnico della teoria dei campi ordinati (ossia l’esistenza dell’estremo superiore per ogni sottoinsieme non vuoto e superiormente limitato), ma lo usa per indicare che, con il sistema dei numeri reali, le estensioni numeriche sono “completate”, dato che, per definizione, non è ulteriormente possibile avere un campo ordinato archimedeo “più grande” del campo reale. La Sezione 2.6.3 mostrerà però che i due significati sono equivalenti, ossia che $\mathbb{R}_H = \mathbb{R}$.

La seconda riguarda il termine “*Erweiterung*”, cioè *estensione*, *ampliamento*, che non va inteso nel senso tecnico della teoria dei campi, ma nel senso maggiormente ristretto che si ha nel caso dei campi ordinati, in particolare archimedei.

Conseguentemente, per poter interpretare correttamente l’*Axiom der Vollständigkeit*, è necessario precisare cosa si intenda per *estensione* sia nel caso dei campi (cf. Sezione 2.6.1) che nel caso dei campi ordinati (cf. Sezione 2.6.2), in particolare archimedei. Le due nozioni sono necessariamente distinte in quanto nel caso dei campi di deve tener conto delle due operazioni addizione e moltiplicazione, mentre nel caso dei campi ordinati occorre tener conto anche della relazione d’ordine.

2.6.1 Estensioni di campi

Richiamiamo alcuni concetti di base della teoria dei campi.

Ziel genau durcharbeitet und in geeigneter Weise modifiziert.

Cf. anche [33], p. 991: *the [real] numbers form a system of objects which cannot be enlarged with the preceding axioms [corpo ordinato archimedeo] continuing to hold.*

Definizione 2.6.1 Sia $(\mathbb{L}, +, \cdot)$ un campo e sia $\mathbb{K}$ un sottoinsieme non vuoto di $\mathbb{L}$. Si dice che $\mathbb{K}$ è un sottocampo di $(\mathbb{L}, +, \cdot)$ se $(\mathbb{K}, +, \cdot)$ è un campo rispetto alle stesse operazioni “+” e “.” definite in $\mathbb{L}$.

Definizione 2.6.2 Si dice che un campo $\mathbb{L}$ è un'estensione di un campo $\mathbb{K}$ se $\mathbb{K}$ è isomorfo ad un sottocampo $\mathbb{K}'$ di $\mathbb{L}$.¹²

Per indicare un'estensione $\mathbb{L}$ di un campo $\mathbb{K}$ si scrive anche $\mathbb{L} | \mathbb{K}$.

Definizione 2.6.3 Se un campo $\mathbb{L}$ è un'estensione di un campo $\mathbb{K}$, l'estensione $\mathbb{L} | \mathbb{K}$ si dice propria se $\mathbb{L}$ non è isomorfo a $\mathbb{K}$.

Osservazione 2.6.4 È importante osservare che, affinché un'estensione $\mathbb{L} | \mathbb{K}$ sia propria non è sufficiente che $\mathbb{K}$ sia isomorfo ad un sottocampo $\mathbb{K}'$ di $\mathbb{L}$ con $\mathbb{K}' \subseteq \mathbb{L}$, $\mathbb{K}' \neq \mathbb{L}$, ma è anche necessario che $(\mathbb{L}, +, \cdot)$ e $(\mathbb{K}, +, \cdot)$ non siano isomorfi. Ad esempio, il campo $\mathbb{R}$ dei numeri reali è un'estensione propria del campo $\mathbb{Q}$ dei numeri razionali perché non solo $\mathbb{Q} \subseteq \mathbb{R}$, $\mathbb{Q} \neq \mathbb{R}$, ma si ha anche che i campi $\mathbb{R}$ e $\mathbb{Q}$ non sono isomorfi.¹³ Analogamente si ha che il campo $\mathbb{C}$ dei numeri complessi è un'estensione propria del campo $\mathbb{R}$ dei numeri reali.

D'altra parte, esistono campi $\mathbb{K}_1$ e $\mathbb{K}_2$ che sono isomorfi pur avendosi $\mathbb{K}_2 \subseteq \mathbb{K}_1$, $\mathbb{K}_2 \neq \mathbb{K}_1$. Ad esempio, si considerino un'infinità numerabile di indeterminate $X_1, X_2, \dots$. Sia $\mathbb{K}_1 = \mathbb{Q}(X_1, X_2, \dots)$ il campo delle funzioni razionali su $\mathbb{Q}$ nelle indeterminate $X_1, X_2, \dots$, e sia $\mathbb{K}_2 = \mathbb{Q}(X_2, X_3, \dots)$ il campo delle funzioni razionali su $\mathbb{Q}$ nelle indeterminate $X_2, X_3, \dots$. Il campo $\mathbb{K}_2$ è un campo contenuto in $\mathbb{K}_1$, ma $\mathbb{K}_1$ non è un'estensione propria di $\mathbb{K}_2$ perché $\mathbb{K}_1$ e $\mathbb{K}_2$ sono isomorfi: infatti esiste un isomorfismo $\phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$ definito da $\phi(r) = r$ per ogni $r \in \mathbb{Q}$, e da $\phi(X_i) = X_{i+1}$ per $i = 1, 2, 3, \dots$.

2.6.2 Estensioni di campi ordinati

Sia $\mathbb{L}$ un campo ordinato, e sia $\mathbb{K}'$ un sottocampo di $\mathbb{L}$. Con l'ordinamento indotto da $\mathbb{L}$, il campo $\mathbb{K}'$ è un campo ordinato.

¹² Se un campo $\mathbb{L}$ è un'estensione di un campo $\mathbb{K}$, allora esiste un omomorfismo $\phi : \mathbb{K} \rightarrow \mathbb{L}$, necessariamente iniettivo: quindi, nel linguaggio della teoria delle categorie, le estensioni sono i morfismi della categoria i cui oggetti sono i campi.

¹³ Se esistesse un isomorfismo di campo $\phi : \mathbb{R} \rightarrow \mathbb{Q}$, allora $\phi(\sqrt{2})$ sarebbe una soluzione razionale dell'equazione $x^2 = 2$.

Definizione 2.6.5 Si dice che un campo ordinato $\mathbb{L}$ è un'estensione ordinata di un campo ordinato $\mathbb{K}$ se $\mathbb{K}$ è ordinatamente isomorfo ad un sottocampo $\mathbb{K}'$ di $\mathbb{L}$.

Definizione 2.6.6 Se un campo ordinato $\mathbb{L}$ è un'estensione ordinata di un campo ordinato $\mathbb{K}$, si dice che l'estensione ordinata $\mathbb{L}|\mathbb{K}$ è propria se non esistono isomorfismi ordinati di $\mathbb{K}$ su $\mathbb{L}$.

Osservazione 2.6.7 È importante osservare che esistono campi ordinati $\mathbb{K}_1$ e $\mathbb{K}_2$ che sono ordinatamente isomorfi pur avendosi $\mathbb{K}_2 \subseteq \mathbb{K}_1$, $\mathbb{K}_2 \neq \mathbb{K}_1$. Basta ad esempio considerare i campi $\mathbb{K}_1 = \mathbb{Q}(X_1, X_2, \dots)$ e $\mathbb{K}_2 = \mathbb{Q}(X_2, X_3, \dots)$ introdotti nell'Osservazione 2.6.4, che diventano campi ordinati se si introduce l'ordinamento indotto da:

$$r < X_1 < X_2 < X_3 < \dots < X_n < \dots$$

dove $r \in \mathbb{Q}$ mentre $X_1, X_2, \dots$ sono le indeterminate (numerabili). Si ha $\mathbb{K}_2 \subseteq \mathbb{K}_1$, $\mathbb{K}_2 \neq \mathbb{K}_1$, ma $\mathbb{K}_1$ non è un'estensione ordinata propria di $\mathbb{K}_2$ perché $\mathbb{K}_1$ e $\mathbb{K}_2$ sono ordinatamente isomorfi tramite lo stesso isomorfismo $\phi: \mathbb{K}_2 \rightarrow \mathbb{K}_1$ definito, come nell'Osservazione 2.6.4, da $\phi(r) = r$ per ogni $r \in \mathbb{Q}$, e da $\phi(X_i) = X_{i-1}$ per $i = 2, 3, \dots$

Caso dei campi ordinati archimedei

Per la nostra trattazione è importante il caso dei campi ordinati *archimedei*, per i quali, come mostra la seguente Proposizione, per stabilire se $\mathbb{L}$ è una estensione ordinata propria di un suo sottocampo $\mathbb{K}'$, non occorre provare che non esistono isomorfismi ordinati di $\mathbb{K}'$ su $\mathbb{L}$, ma è sufficiente verificare la semplice inclusione insiemistica stretta, ossia che $\mathbb{K}' \subseteq \mathbb{L}$ con $\mathbb{K}' \neq \mathbb{L}$.

Proposizione 2.6.8 Siano $\mathbb{L}$ e $\mathbb{K}$ campi ordinati archimedei. Il campo $\mathbb{L}$ è un'estensione ordinata propria di $\mathbb{K}$ se e solo se $\mathbb{K}$ è ordinatamente isomorfo ad un sottocampo $\mathbb{K}'$ di $\mathbb{L}$ con $\mathbb{K}' \subseteq \mathbb{L}$ con $\mathbb{K}' \neq \mathbb{L}$.

DIMOSTRAZIONE. Supponiamo che $\mathbb{L}$ sia una estensione ordinata *propria* di $\mathbb{K}$: dalla definizione di estensione ordinata e di estensione ordinata propria segue banalmente che $\mathbb{K}$ è isomorfo ad un campo $\mathbb{K}'$ con $\mathbb{K}' \subseteq \mathbb{L}$, $\mathbb{K}' \neq \mathbb{L}$.

Viceversa, supponiamo che esista un isomorfismo ordinato $\sigma : \mathbb{K} \rightarrow \mathbb{K}'$ con $\mathbb{K}'$ sottocampo di $\mathbb{L}$, $\mathbb{K}' \neq \mathbb{L}$. Tenuto conto di σ , per provare che non esistono isomorfismi ordinati $\Phi : \mathbb{K} \rightarrow \mathbb{L}$, ossia che l'estensione $\mathbb{L}$ è propria, è sufficiente provare che non esistono isomorfismi ordinati fra $\mathbb{K}'$ ed $\mathbb{L}$. Supponiamo per assurdo che esista un isomorfismo ordinato $\Psi : \mathbb{K}' \rightarrow \mathbb{L}$; sia $y \in \mathbb{L} \setminus \mathbb{K}'$ e sia $x = \Psi^{-1}(y) \in \mathbb{K}'$. Poiché $x \in \mathbb{K}'$ e $\Psi(x) = y \notin \mathbb{K}'$, si ha $x \neq \Psi(x)$; supponiamo sia $x < \Psi(x)$. L'ordinamento di $\mathbb{L}$ è archimedeo e pertanto, per la Proposizione 2.3.4, esiste un razionale r tale che $x < r < \Psi(x)$, inoltre $\Psi(r) = r$ (cf. (2.3)) e pertanto si ha $\Psi(x) < \Psi(r) = r$ che contraddice la disuguaglianza precedente. Analogamente per il caso $\Psi(x) < x$. $\square$

Conseguentemente, se un campo ordinato archimedeo $\mathbb{L}$ è un'estensione ordinata propria di un campo ordinato archimedeo $\mathbb{K}$, è sempre possibile identificare $\mathbb{K}$, a meno di isomorfismi ordinati, con un sottocampo di $\mathbb{L}$ diverso da $\mathbb{L}$. Possiamo ora precisare formalmente l'*Axiom der Vollständigkeit* di Hilbert:

Definizione 2.6.9 *Un campo ordinato archimedeo $\mathbb{K}$ verifica l'Axiom der Vollständigkeit se $\mathbb{K}$ non ha estensioni ordinate proprie che siano archimedee.*

Definizione 2.6.10 *In modo assiomatico si definisce $\mathbb{R}_H$ (campo dei numeri reali “secondo Hilbert”) come un campo ordinato archimedeo che verifica l'Axiom der Vollständigkeit, ossia come un campo ordinato archimedeo che non ha estensioni ordinate proprie archimedee.*

2.6.3 La *Vollständigkeit* equivale alla completezza

Possiamo ora dimostrare l'equivalenza della Definizione 2.6.10 del sistema dei numeri reali data da Hilbert con quella “moderna” espressa dalla Definizione 2.5.9. Iniziamo con il dimostrare che:

Proposizione 2.6.11 *Se $\mathbb{K}$ è un campo ordinato che è ordinatamente isomorfo ad un campo completo, allora $\mathbb{K}$ è completo.*

DIMOSTRAZIONE. Sia Φ un isomorfismo ordinato fra un campo ordinato $\mathbb{K}$ ed un campo ordinato completo $\mathbb{R}$. Sia dato un sottoinsieme $E \subseteq \mathbb{K}$ non-vuoto e superiormente limitato, ad esempio da un elemento $m \in \mathbb{K}$.

Allora $\Phi(E)$ è un sottoinsieme di $\mathbb{R}$ che non è vuoto (perché E non è vuoto) ed è superiormente limitato da $\Phi(m)$ (in quanto da $x \leq m$ segue $\Phi(x) \leq \Phi(m)$). Esiste quindi $\xi = \sup \Phi(E)$, e valgono le due proprietà caratteristiche (vedi Osservazione 2.4.2):

- (i) Per ogni $x \in \Phi(E)$ si ha $x \leq \xi$.
- (ii) Per ogni $\epsilon > 0$ esiste $\bar{x} \in \Phi(E)$ tale che $\bar{x} < \xi - \epsilon$.

Applicando l'isomorfismo ordinato inverso Φ^{-1} , che conserva le operazioni e l'ordine, e ponendo $a = \Phi^{-1}(x)$, $\delta = \Phi^{-1}(\epsilon)$, $\bar{a} = \Phi^{-1}(\bar{x})$, si ha:

- (i) Per ogni $a \in E$ si ha $a \leq \Phi^{-1}(\xi)$.
- (ii) Per ogni $\delta > 0$ esiste $\bar{a} \in E$ tale che $\bar{a} < \Phi^{-1}(\xi) - \delta$.

Conseguentemente, $\Phi^{-1}(\xi)$ è l'estremo superiore di E , ed il campo ordinato $\mathbb{K}$ è completo. $\square$

Teorema 2.6.12 *Un campo ordinato archimedeo soddisfa l'Axiom der Vollständigkeit di Hilbert se e solo se è un campo ordinato completo.*

DIMOSTRAZIONE. Sia $\mathbb{K}$ un campo ordinato completo, e dunque anche archimedeo per la Proposizione 2.4.5; dimostriamo che $\mathbb{K}$ soddisfa l'*Axiom der Vollständigkeit*. Dobbiamo quindi provare, data un'estensione ordinata archimedeo $\mathbb{L}$ di $\mathbb{K}$, questa non può essere propria. Supponiamo quindi che esista un isomorfismo ordinato Ψ di $\mathbb{K}$ su un sottocampo $\mathbb{K}' \subseteq \mathbb{L}$. Conseguentemente, per la Proposizione 2.6.11, il campo $\mathbb{K}'$ è completo, e quindi, per il Teorema di Hilbert 2.5.3, esiste un omomorfismo ordinato $\Phi : \mathbb{L} \rightarrow \mathbb{K}'$. Il sottocampo $\Phi(\mathbb{L}) \subseteq \mathbb{K}'$ è ordinatamente isomorfo ad $\mathbb{L}$. Abbiamo quindi $\Phi(\mathbb{L}) \subseteq \mathbb{K}' \subseteq \mathbb{L}$: nessuna delle due inclusioni può essere stretta, in quanto altrimenti avremmo un sottocampo proprio $\Phi(\mathbb{L})$ del un campo archimedeo $\mathbb{L}$ che risulterebbe ordinatamente isomorfo ad $\mathbb{L}$ stesso, contro la Proposizione 2.6.8. Quindi è $\Phi(\mathbb{L}) = \mathbb{K}' = \mathbb{L}$, e pertanto $\mathbb{K}$ è ordinatamente isomorfo (tramite l'isomorfismo ordinato Ψ) ad $\mathbb{L} = \mathbb{K}'$.

L'implicazione "*Axiom der Vollständigkeit* $\Rightarrow$ *completo*" può essere provata utilizzando un modello di campo ordinato completo, come ad esempio quello $\mathbb{R}$ di Méray–Cantor o quello di Dedekind. Infatti, sia $\mathbb{K}$ un campo ordinato archimedeo che soddisfa l'*Axiom der Vollständigkeit*. Per il Teorema di Hilbert 2.5.3, $\mathbb{K}$ è ordinatamente isomorfo (tramite un isomorfismo ordinato Φ) ad un sottocampo $\Phi(\mathbb{K})$ di $\mathbb{R}$. Per la Proposizione 2.4.5, $\mathbb{R}$ è archimedeo, ed è ovvio che anche $\Phi(\mathbb{K})$ soddisfa l'*Axiom der*

Vollständigkeit. Quindi $\Phi(\mathbb{K})$ non può essere un sottocampo proprio di $\mathbb{R}$, ed è $\Phi(\mathbb{K}) = \mathbb{R}$. In conclusione, $\mathbb{K}$ è ordinatamente isomorfo (tramite l'isomorfismo ordinato Φ) al campo completo $\Phi(\mathbb{K}) = \mathbb{R}$, e pertanto, per la Proposizione 2.6.11, $\mathbb{K}$ è un campo ordinato completo.

Alternativamente, per dimostrare che un campo ordinato archimedeo $\mathbb{K}$ che soddisfa l'*Axiom der Vollständigkeit* è necessariamente completo può essere utilizzato un risultato di B. Banaschewski [3], che afferma che se $\mathbb{K}$ è un campo ordinato archimedeo che non è completo (e pertanto esiste un sottoinsieme $S \subseteq \mathbb{K}$ non-vuoto, superiormente limitato, ma privo di estremo superiore) è possibile costruire esplicitamente (in dipendenza da S) un'estensione ordinata propria archimedeo di $\mathbb{K}$, contro l'*Axiom der Vollständigkeit*. $\square$

Osservazione 2.6.13 Si osservi che nel Teorema 2.6.12 l'implicazione *completo* $\implies$ *Axiom der Vollständigkeit* è provata a prescindere dall'esistenza di un campo ordinato completo. Per l'implicazione opposta *Axiom der Vollständigkeit* $\implies$ *completo*, abbiamo fornito due dimostrazioni diverse. La prima sfrutta esplicitamente l'esistenza di un modello di $\mathbb{R}$. La seconda invece, basata sul risultato di [3], non ha bisogno dell'esistenza *a priori* di un campo ordinato completo. Resterebbe quindi irrisolto il problema di capire in quale modo David Hilbert intendesse dimostrare il fondamentale teorema di esistenza dell'estremo superiore per il “suo” sistema dei numeri reali.

2.6.4 “Esistenza”: un'idea per un modello di $\mathbb{R}$

Abbiamo visto come l'*Axiom der Vollständigkeit* di Hilbert definisca $\mathbb{R}$ come un campo ordinato archimedeo *massimale* (nel senso precisato dalla Proposizione 2.6.8). In questa sezione mostriamo come questa definizione “astratta” contenga, in modo molto naturale, l'idea per una costruzione “concreta” di un modello di $\mathbb{R}$ (precisamente per il modello delle “se-mirette di Russell”, una variante del modello di Dedekind presentato nel Capitolo 4).

Il primo passo è la constatazione che ogni elemento di un campo ordinato archimedeo $\mathbb{K}$ è univocamente determinato dagli elementi razionali che lo precedono.

Proposizione 2.6.14 *Sia $\mathbb{K}$ un campo ordinato archimedeo. Sia $x \in \mathbb{K}$, e sia $L(x) = \{r \in \mathbb{Q} \mid r < x\}$. Allora $x = \sup L(x)$.*

DIMOSTRAZIONE. Per definizione, x è un maggiorante di $L(x)$. Se y fosse una maggiorazione di $L(x)$ minore di x , esisterebbe, per la Proposizione 2.3.4, un razionale r tale che $y < r < x$. Sarebbe allora $r \in L(x)$, contraddicendo l'assunzione che y sia una maggiorazione di $L(x)$. $\square$

Notiamo che quando $x > 0$, ponendo $L^+(x) = \{r \in \mathbb{Q} \mid 0 < r < x\}$, si ha ancora che $x = \sup L^+(x)$. Definiamo ora le *semirette di Russell*.

Definizione 2.6.15 *Un sottoinsieme $A \subseteq \mathbb{Q}$ è detto una semiretta di Russell se è*

- (R1') *superiormente limitato,*
- (R2') *non-vuoto,*
- (R3') *privo di massimo,*
- (R4') *inferiormente saturo; ossia, per ogni $a' \in \mathbb{Q}$, se $a' < a \in A$ si ha che $a' \in A$.*

Proposizione 2.6.16 *Sia $\mathbb{K}$ un campo ordinato archimedeo e sia $\mathcal{R}$ l'insieme delle semirette di Russell.¹⁴ Per ogni $x \in \mathbb{K}$ il sottoinsieme $L(x)$ è una semiretta di Russell, e l'applicazione $L : \mathbb{K} \rightarrow \mathcal{R}$ è strettamente crescente (rispetto all'ordinamento dato in $\mathbb{K}$ ed all'ordinamento per inclusione in $\mathcal{R}$).*

DIMOSTRAZIONE. Sia dato $x \in \mathbb{K}$. (R1'): $L(x)$ è superiormente limitato da x . (R2'): essendo $\mathbb{K}$ archimedeo, esiste un naturale n tale che $n \cdot 1 > -x$, ossia $-n \cdot 1 < x$, e quindi $-n \in L(x)$. (R3'): supponiamo per assurdo che $L(x)$ abbia un massimo elemento m ; da $m \in L(x)$ segue $m < x$ e dalla Proposizione 2.3.4 segue che esiste un razionale r tale che $m < r < x$; avendosi quindi $r \in L(x)$ si ha un assurdo. (R4'): se $r' < r \in L(x)$ si ha $r' < r < x$, dal che $r' \in L(x)$, ossia $L(x)$ è inferiormente saturo.

Siano ora $x, y \in \mathbb{K}$, $x \neq y$. Supponiamo ad esempio che sia $x < y$. Ogni razionale minore di x è anche minore di y , da cui $L(x) \subseteq L(y)$. Per la Proposizione 2.3.4 esiste un razionale r tale che $x < r < y$. Si ha quindi $r \in L(y)$, $r \notin L(x)$, ossia $L(x) \neq L(y)$. $\square$

La Proposizione precedente mostra in sostanza che ogni elemento di un campo ordinato archimedeo $\mathbb{K}$ è identificabile con una semiretta di Russell. Inoltre, dato che fra semirette di Russell possono essere definiti in

¹⁴ Si noti che $\mathcal{R}$ è un sottoinsieme dell'insieme $\mathcal{P}(\mathbb{Q})$ delle parti di $\mathbb{Q}$. Questo pone una limitazione alla cardinalità di un campo ordinato archimedeo, che non può superare quella dell'insieme $\mathcal{P}(\mathbb{Q})$.

maniera autonoma, cioè indipendente da $\mathbb{K}$, le operazioni di addizione e moltiplicazione e la relazione di ordine,¹⁵ è ragionevole pensare che, utilizzando la Proposizione 2.6.14, sia possibile costruire un isomorfismo ordinato fra $\mathbb{K}$ ed un *sottoinsieme* di semirette di Russell strutturato a campo ordinato con le operazioni dette.

L'idea fondamentale per un modello di $\mathbb{R}$, visto l'*Axiom der Vollständigkeit*, è allora quella di considerare l'insieme $\mathcal{R}$ di *tutte* le semirette di Russell, strutturato a campo ordinato con le operazioni dette. Un tale campo infatti, per la Proposizione 2.6.8, non potrà avere estensioni ordinate archimedee proprie, e quindi avremo in sostanza, simbolicamente, che $\mathcal{R} = \mathbb{R}_H = \mathbb{R}$.

Questa è infatti l'idea del modello di $\mathbb{R}$ di B. Russell [58], che verrà discusso in dettaglio nel Capitolo 4. In effetti il Capitolo 4 è dedicato al più noto e classico modello di Dedekind, del quale il modello di Russell è una “semplificazione”.

2.6.5 Una precisazione terminologica

In una certa tradizione, soprattutto di ambiente algebrico, è piuttosto frequente incontrare l'espressione “campo ordinato archimedeo completo” per indicare il campo $\mathbb{R}$ dei numeri reali. È ragionevole pensare che questa espressione sia ereditata dall'assiomatica hilbertiana, in cui, come si è visto, $\mathbb{R}$ viene individuato come un “campo ordinato” con l'aggiunta dell'assioma di Archimede (“archimedeo”) e dell'*Axiom der Vollständigkeit* (“completo”, dove questo aggettivo va qui inteso, ovviamente, nel senso di Hilbert: il campo non ammette estensioni “in cui continuo a valere tutti gli altri assiomi”). Se invece, nell'espressione “campo ordinato archimedeo completo” l'aggettivo “completo” ha il significato della Definizione 2.4.1 (si riveda a questo proposito l'Osservazione 2.4.3), allora, per la Proposizione 2.4.5, l'aggettivo “archimedeo” risulta superfluo.

¹⁵ La “somma” di due semirette A, B è semplicemente $A + B$, mentre la definizione del “prodotto” $A \cdot B$ di due semirette A, B è analoga, pur richiedendo qualche complicazione dovuta alla necessità di distinguere i segni. L’“ordine” in $\mathcal{R}$ è definito da $A \leq B$ se e solo se $A \subseteq B$.

2.7 Proprietà di Dedekind e di Cantor

Alcune impostazioni dell'analisi matematica preferiscono definire assiomaticamente $\mathbb{R}$ non come un campo ordinato completo (unico a meno di isomorfismi ordinati), ma come un campo ordinato nel quale valgono opportune proprietà di “separazione”.

Definizione 2.7.1 *Si dice che una coppia (A, B) di sottoinsiemi non vuoti di un campo ordinato $\mathbb{K}$ è una coppia di classi separate se per ogni $a \in A$ ed ogni $b \in B$ si ha $a \leq b$.*

Si dice che una coppia (A, B) di sottoinsiemi non vuoti di un campo ordinato $\mathbb{K}$ è una coppia di classi contigue se (i) è una coppia di classi separate e (ii) per ogni $\epsilon \in \mathbb{K}$, $\epsilon > 0$ esiste $a \in A$ e $b \in B$ tali che $b - a < \epsilon$.

Si dice che una coppia (A, B) di classi separate ammette un elemento separatore se esiste un elemento $x \in \mathbb{K}$ tale che per ogni $a \in A$ ed ogni $b \in B$ si ha $a \leq x \leq b$.

Si dice che in un campo ordinato $\mathbb{K}$ vale la Proprietà di Dedekind se ogni coppia (A, B) di classi separate ammette un elemento separatore.

Si dice che in un campo ordinato $\mathbb{K}$ vale la Proprietà di Cantor se ogni coppia (A, B) di classi contigue ammette un elemento separatore.

Evidentemente in un campo ordinato la Proprietà di Dedekind implica la Proprietà di Cantor, in quanto le classi contigue sono *particolari* classi separate. Intuitivamente, si potrebbe credere vero anche il viceversa, ossia che se ci sono elementi separatori fra classi contigue, che sembrano “incollate” assieme, con “poco spazio” a disposizione, a maggior ragione ce ne saranno fra le classi separate, fra le quali il “buco” potrebbe sembrare, caso mai, più largo. In effetti le due proprietà non sono equivalenti.

Convienne, anche didatticamente, esaminare criticamente la classica dimostrazione dell'analisi matematica di esistenza dell'estremo superiore basata sulla Proprietà di Dedekind.

Teorema 2.7.2 *Sia $\mathbb{K}$ un campo ordinato in cui vale la Proprietà di Dedekind, e sia $E \subseteq \mathbb{K}$ un sottoinsieme non vuoto e limitato superiormente. Allora l'insieme B delle maggiorazioni di E ha minimo.*

DIMOSTRAZIONE. Sia A il complementare di B . Mostriamo che A e B formano una coppia (A, B) di classi separate. Dati $a \in A$ e $b \in B$, per definizione a non è una maggiorazione di E : esiste quindi un elemento $u \in E$ tale che $u > a$; essendo b una maggiorazione di E , risulta $b \geq u$, e per la proprietà transitiva della relazione d'ordine risulta $a < b$. Per la Proprietà di Dedekind, esiste un elemento separatore ξ , ossia un elemento tale che, per ogni $a \in A$ ed ogni $b \in B$, risulta $a \leq \xi \leq b$. Pertanto, per provare che ξ è il minimo elemento di B , basta dimostrare che $\xi \in B$. Se, per assurdo, fosse $\xi \in A$, cioè se ξ non fosse una maggiorazione di E , esisterebbe un $a' \in E$ tale che $\xi < a'$. Conseguentemente, gli elementi di $\mathbb{K}$ strettamente compresi fra ξ ed $a' \in E$ non sono maggiorazioni di E , ossia appartengono ad A ; scelto uno qualsiasi di questi elementi, ad esempio $a = (\xi + a')/2$ (come in [22], p. 16), oppure un numero razionale a dell'intervallo aperto $]\xi, a'[,$ (cf. Proposizione 2.3.4), si ha una contraddizione con il fatto che per ogni $a \in A$ risulta $a \leq \xi$. $\square$

Mostreteremo ora che, in generale, sostituendo nell'ipotesi del Teorema 2.7.2 la Proprietà di Dedekind con la più debole Proprietà di Cantor, la tesi può ancora essere dimostrata pur di considerare campi ordinati *archimedei*.

Teorema 2.7.3 *Sia $\mathbb{K}$ un campo ordinato archimedeo in cui vale la Proprietà di Cantor, e sia $E \subseteq \mathbb{K}$ un sottoinsieme non vuoto e limitato superiormente. Allora l'insieme B delle maggiorazioni di E ha minimo.*

DIMOSTRAZIONE. Sia A il complementare di B . Mostriamo che A e B formano una coppia (A, B) di classi contigue. Possiamo provare, come nel Teorema 2.7.2, che A e B formano una coppia (A, B) di classi separate. Consideriamo un elemento $a_0 \in A$ ed un elemento $b_0 \in B$.¹⁶ Procediamo per bisezione.

- Consideriamo il punto di mezzo $t_0 = (a_0 + b_0)/2$: definiamo un intervallo

$$[a_1, b_1] = \begin{cases} [a_0, t_0] & \text{se } t_0 \in B \\ [t_0, b_0] & \text{se } t_0 \in A \end{cases}$$

¹⁶ Notiamo che B non è vuoto perché E è limitato superiormente, ovvero esistono maggiorazioni di E ; dato che E non è vuoto, possiamo considerare un elemento $x \in E$, ed allora $x - 1$ non è una maggiorazione di E , ovvero A non è vuoto.

- Consideriamo il punto di mezzo $t_1 = (a_1 + b_1)/2$: definiamo un intervallo

$$[a_2, b_2] = \begin{cases} [a_1, t_1] & \text{se } t_1 \in B \\ [t_1, b_1] & \text{se } t_1 \in A \end{cases}$$

- ...

L'algoritmo definisce, per ogni intero $n \geq 1$, un elemento $a_n \in A$ ed un elemento $b_n \in B$ tali che $b_n - a_n = (b_0 - a_0)/2^n$. Dato un qualsiasi $\epsilon \in \mathbb{K}$, $\epsilon > 0$, per la Proprietà di Archimede esiste un intero $m > (b_0 - a_0)\epsilon^{-1}$. Scelta una potenza 2^n di 2 maggiore di m , risulta $(b_0 - a_0)/2^n < \epsilon$, e quindi $|b_n - a_n| = b_n - a_n = (b_0 - a_0)/2^n < \epsilon$. Quindi A e B formano una coppia (A, B) di classi contigue. Per la Proprietà di Cantor, esiste un elemento separatore ξ e la dimostrazione può essere conclusa come nel Teorema 2.7.2. $\square$

Osserviamo che nella dimostrazione del Teorema 2.7.3 il ruolo della Proprietà di Archimede è essenziale. Infatti l' $\epsilon > 0$ nella definizione di classi contigue è del tutto arbitrario: se non valesse la Proprietà di Archimede, ϵ potrebbe essere un infinitesimo, e in tal caso, per quanto grande sia n , potrebbe essere sempre $0 < \epsilon \leq (b_0 - a_0)/2^n$. Ad esempio, nel campo non archimedeo $\mathbb{F}$ dell'Esempio 2.3.3, l'insieme $E = \{-\frac{1}{n} \mid n \geq 1\}$, pur essendo non vuoto e superiormente limitato (da 0), non ha estremo superiore. Ad esempio, 0 non è la minima maggiorazione, in quanto $b = -\frac{1}{x}$ è negativo ed è una maggiorazione di E . Infatti per ogni $n \geq 1$ risulta

$$-\frac{1}{n} < -\frac{1}{x} \quad \text{poiché} \quad -\frac{1}{x} - \left(-\frac{1}{n}\right) = \frac{x - n}{nx} > 0$$

Proposizione 2.7.4 *In un campo ordinato vale la Proprietà di Dedekind se e solo se il campo è completo.*

DIMOSTRAZIONE. È stato dimostrato nel Teorema 2.7.2 che in un campo ordinato la Proprietà di Dedekind implica che il campo è completo. Per provare che in un campo completo vale la Proprietà di Dedekind, basta notare che se (A, B) sono separate, allora $\sup A$ è un elemento di separazione. $\square$

Quanto dimostrato rende possibile introdurre assiomaticamente $\mathbb{R}$ come un campo ordinato in cui valgono opportune proprietà di “separazione”, precisamente:

- (a) come un campo ordinato nel quale vale la Proprietà di Dedekind, oppure equivalentemente, come
- (b) come un campo ordinato *archimedeo* in cui vale la Proprietà di Cantor.

La Tabella 2.7 rappresenta un riassunto schematico dei legami fra la Proprietà di Cantor, la Proprietà di Archimede, la Proprietà di Dedekind e la completezza nei campi ordinati.

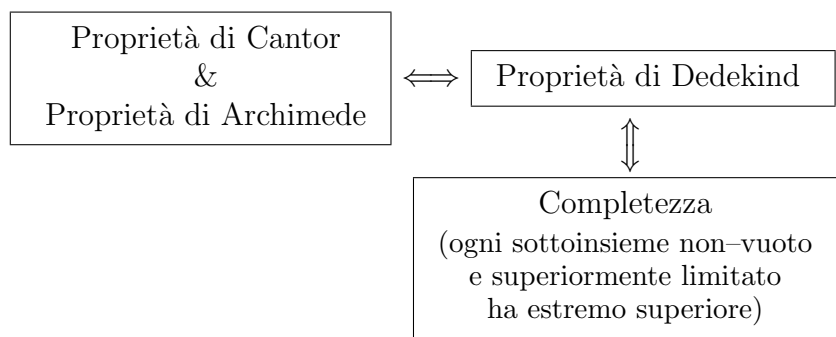


Tabella 2.1: Legami fra la Proprietà di Cantor, la Proprietà di Archimede, la Proprietà di Dedekind e la completezza nei campi ordinati.

Naturalmente, nel caso di una definizione assiomatica di $\mathbb{R}$ basata su proprietà di “separazione”, diventa poi necessario dimostrare il teorema di esistenza dell’estremo superiore, ciò che non accade nel caso della definizione assiomatica di $\mathbb{R}$ come campo ordinato completo (Definizione 2.5.9), dato che, in tal caso, ogni sottoinsieme di $\mathbb{R}$ non vuoto e superiormente limitato ha l’estremo superiore *per definizione*. Sui diversi percorsi introduttivi dell’analisi matematica si veda il successivo Capitolo 5.

Osservazione 2.7.5 Taluni testi, ad esempio [22], adottano una definizione alternativa della Proprietà di Dedekind, formalmente diversa da quella della Definizione 2.7.1, richiedendo l’esistenza di un elemento separatore per le coppie (A, B) di classi separate *che soddisfano l’ulteriore condizione di essere una partizione di $\mathbb{K}$* .

Le due definizioni sono equivalenti. Infatti, se in un campo ordinato $\mathbb{K}$ ammettono elemento separatore le coppie (A, B) di classi separate che sono una

partizione di $\mathbb{K}$, allora, per la dimostrazione del Teorema 2.7.2, il campo $\mathbb{K}$ è completo, e per la Proposizione 2.7.4 vale in $\mathbb{K}$ la Proprietà di Dedekind della Definizione 2.7.1.

Didatticamente conviene enfatizzare il fatto, piuttosto controintuitivo, che la summenzionata “definizione alternativa” della Proprietà di Dedekind, malgrado essa richieda che le due classi “esauriscano” tutto il campo, non definisce classi contigue. Infatti:

Proposizione 2.7.6 *Esistono classi separate che sono una partizione di un campo ordinato $\mathbb{K}$, ma non sono contigue.*

DIMOSTRAZIONE. Consideriamo nuovamente il campo non archimedeo $\mathbb{F}$ dell'Esempio 2.3.3. Sia $E = \{-\frac{1}{n} \mid n \geq 1\}$, sia B l'insieme delle maggiorazioni di E e sia $A = \mathbb{F} \setminus B$. Ragionando come nella dimostrazione del Teorema 2.7.2 otteniamo che A e B sono separate e costituiscono ovviamente una partizione del campo $\mathbb{F}$. Proviamo che A e B non sono contigue. Fissiamo un $\epsilon^* > 0$ infinitesimo, e consideriamo elementi arbitrari $a \in A$ e $b \in B$. Si ha $b - a \geq \epsilon^*$. Infatti, se fosse $b - a < \epsilon^*$, scelto un elemento $-1/n$ di E tale che $-1/n > a$, ossia $1/n < -a$, si avrebbe $b + 1/n < b - a < \epsilon^*$, e questo è assurdo, sia quando $b \geq 0$ che quando $b < 0$, perché in questo caso b , essendo una maggiorazione di E , è necessariamente un infinitesimo.

Un altro esempio nello stesso campo: sia $E = \mathbb{N}$, sia B l'insieme delle maggiorazioni di E e sia $A = \mathbb{F} \setminus B$. Le classi A e B sono separate e sono una partizione di $\mathbb{F}$. Consideriamo elementi arbitrari $a \in A$ e $b \in B$. Scelto $n \in \mathbb{N}$ tale che $a < n$, si ha $a < n < n + 1 < b$, per cui $b - a > 1 + n - a > 1$. In generale, se $a \in A$ e $b \in B$, quale che sia l'intero $m \geq 1$ si ha $b - a > m$. $\square$

Chiudiamo la sezione osservando che:

Osservazione 2.7.7 *Anche i reali, come i razionali, non posseggono automorfismi diversi dall'identità.*

Infatti se $f : \mathbb{R} \rightarrow \mathbb{R}$, è un automorfismo,¹⁷ per ogni x, y si ha $f(x + y) = f(x) + f(y)$, ed $f(xy) = f(x)f(y)$; allora risulta $f(r) = r$ per $r \in \mathbb{Q}$; da $f(x + h^2) = f(x) + f(h^2) = f(x) + f(h)^2$ si deduce che $x < y = x + h^2$ implica $f(x) < f(y)$; se ora $r < x < s$ con r, s razionali, si ha $r < f(x) < s$. Per la arbitrarietà di r, s si ha $f(x) = x$ per ogni $x \in \mathbb{R}$. $\square$

A differenza di $\mathbb{Q}$ e di $\mathbb{R}$, il campo $\mathbb{C}$ dei numeri complessi ha due automorfismi che tengono fissi i numeri reali, l'automorfismo identità

¹⁷ Un automorfismo di un campo è un isomorfismo del campo in sé stesso.

$a + ib \mapsto a + ib$ e l'automorfismo coniugio che ad ogni numero complesso $a + ib$ associa il suo coniugato $a - ib$ (ossia le radici quadrate di -1 sono indistinguibili); inoltre esistono infiniti automorfismi “selvaggi” del campo $\mathbb{C}$ dei numeri complessi, costruiti via Lemma di Zorn (vedi [59]), che non tengono fissi tutti i reali.

2.8 Commenti

Nei prossimi due capitoli illustreremo le costruzioni di due modelli per $\mathbb{R}$: quello di Cantor–Méry e quello di Dedekind. Entrambe le costruzioni sono sostanzialmente delle procedure di estensione che, partendo da un campo ordinato archimedeo $\mathbb{K}$, lo estendono ad un campo ordinato $\mathbb{F}$. Sia la procedura di Cantor–Méry che la procedura di Dedekind partono dal campo $\mathbb{Q}$ dei razionali e lo estendono al campo $\mathbb{R}$ dei reali. Entrambe le costruzioni hanno un loro teorema fondamentale, che asserisce sostanzialmente la “completezza hilbertiana”, nel senso che, in entrambe, se si partisse dal campo $\mathbb{R}$ si ritroverebbe come estensione lo stesso $\mathbb{R}$.

Come si vedrà nel Capitolo 3, la costruzione di Cantor–Méry considera le *successioni di Cauchy* in $\mathbb{Q}$; le classifica (modulo un'opportuna equivalenza) in due specie: quelle della prima specie sono quelle convergenti in $\mathbb{Q}$, e vengono identificate con i loro limiti, i razionali; quelle della seconda specie sono quelle non convergenti in $\mathbb{Q}$, e vengono identificate con i “nuovi numeri”, gli irrazionali. Se si applicasse la stessa costruzione ad $\mathbb{R}$ non si produrrebbero “nuovi numeri”, in quanto il teorema fondamentale della costruzione di Cantor–Méry afferma, come si vedrà, che in $\mathbb{R}$ tutte le successioni di Cauchy sono convergenti, ossia in $\mathbb{R}$ ci sono solo successioni di Cauchy di prima specie.

Come illustrato nel Capitolo 4, la costruzione di Dedekind prende in considerazione le cosiddette *sezioni* di $\mathbb{Q}$; le classifica in due specie: quelle della prima specie sono sostanzialmente create da un numero “sezionatore” razionale con il quale vengono identificate; quelle della seconda specie sono tutte le altre, che vengono identificate con i “nuovi numeri”, gli irrazionali. Anche in questo caso, se la stessa costruzione fosse applicata ad $\mathbb{R}$ non si produrrebbero “nuovi numeri”, in quanto il teorema fondamentale della costruzione di Dedekind afferma, come si vedrà, che in $\mathbb{R}$ tutte le sezioni sono di prima specie.

I due teoremi fondamentali quindi hanno una stretta parentela fra loro; sono entrambi la realizzazione nei rispettivi modelli dell'*Axiom der Vollständigkeit* di Hilbert.

Capitolo 3

La costruzione di Méray–Cantor

Das Wesen der Mathematik liegt gerade in ihrer Freiheit.^a

GEORG CANTOR

^a L'essenza della matematica è la sua libertà.



(a)



(b)

Figura 3.1: (a) H. Méray (1835–1911) e (b) G. Cantor (1845–1918).

3.1 Introduzione

Nel 1869 Hugues (Charles) Méray pubblica sulla *Revue des Sociétés Savantes* una memoria intitolata *Remarques sur la nature des quantités définies par la condition de servir de limites à des variables données* contenente una teoria aritmetica dei numeri reali. Ovviamente, data l'epoca, il lavoro di Méray non soddisfa gli attuali standard di “rigore” matematico, ma è sufficientemente preciso da poter essere considerato una delle prime se non la prima costruzione pubblicata di $\mathbb{R}$. Méray utilizza nella sua costruzione di $\mathbb{R}$ certi oggetti detti *varianti*, che giocano un ruolo analogo alle successioni di Cauchy di razionali. Successivamente, nel 1872, Méray ripropone la sua costruzione di $\mathbb{R}$ nel trattato di analisi infinitesimale [42], ma tale contributo, fondamentale, viene in seguito quasi “dimenticato”, forse perché il recensore del trattato, Hermann Laurent,¹ lo trascura completamente, ritenendolo eccessivamente pedante per un libro ritenuto di testo.

Fatto è che sempre nello stesso anno, il 1872, Georg Cantor pubblica l'articolo *Über die Ausdehnung eines Satzes aus der Theorie der trigonometrischen Reihen* [11], ossia “sull'estensione di un teorema della teoria delle serie trigonometriche”, nel quale si trova a considerare insiemi infiniti di punti in relazione al problema della convergenza delle serie trigonometriche;² per poter operare rigorosamente, premette una teoria aritmetica dei numeri reali in cui ogni numero reale è rappresentato da una successione di Cauchy (che lui chiama *successione fondamentale*) di numeri razionali. Come vedremo, due successioni di Cauchy di numeri razionali rappresentano lo stesso numero reale se la differenza delle due successioni tende a 0. Sostanzialmente la costruzione di Cantor è la stessa costruzione di Méray.

La costruzione di Méray–Cantor privilegia la struttura algebrica rispetto a quella d'ordine, ed è più vicina ad una visione dinamica di

¹ Hermann Laurent (1841–1908). Non va confuso con Pierre Laurent (1813–1854). Entrambi lavorarono in teoria delle funzioni; la *serie di Laurent* prende il nome da Pierre.

² David R. Wilkins ha riscritto in L^AT_EX questo famoso articolo di Cantor, e lo ha pubblicato nel web all'url: www.maths.tcd.ie/pub/HistMath/People/Cantor/Ausdehnung/. Chi legge e capisce il tedesco può darci utilmente un'occhiata.

“approssimazione” dei reali coi razionali di quanto lo sia la costruzione “statica” di Dedekind che vedremo nel prossimo capitolo.

Nel presente capitolo presentiamo una versione della costruzione di Méray–Cantor, utilizzando qualche conoscenza di base dei corsi universitari di algebra, in particolare sugli anelli quoziente (cf. [18], p. 221), fino a provare la *completezza metrica* di $\mathbb{R}$ (ogni successione reale di Cauchy converge). La costruzione può essere utilizzata anche in un corso di algebra come esempio non banale di applicazione del teorema sul quoziente di un anello commutativo unitario su un ideale massimale. La sezione conclusiva del capitolo studia il problema didattico (e logico) della definizione classica di successione di Cauchy (con quattro quantificatori), proponendo la definizione alternativa di Bolzano (con tre quantificatori) e la definizione basata sulla finitezza (con un solo quantificatore) utilizzata in alcuni verificatori automatici. La Sezione 3.3 è dovuta ad una collaborazione con Antonella Montone; in essa si discutono le problematiche didattiche del concetto di successione di Cauchy, segnalando anche un’errata traduzione o interpretazione, presente in taluni autori, della definizione originale (di successione di Cauchy) data da Cantor.

3.2 La costruzione di Méray–Cantor

Sia $\mathbb{Q}^{\mathbb{N}}$ l’insieme delle successioni $p = (p_0, p_1, p_2, \dots) = (p_k)_{k \geq 0}$ di numeri razionali. In $\mathbb{Q}^{\mathbb{N}}$ si definiscono l’addizione e la moltiplicazione (per componenti): dati $p = (p_0, p_1, p_2, \dots) = (p_k)_{k \geq 0}$ e $q = (q_0, q_1, q_2, \dots) = (q_k)_{k \geq 0}$ in $\mathbb{Q}^{\mathbb{N}}$ poniamo

$$\begin{aligned} p + q &= (p_0 + q_0, p_1 + q_1, p_2 + q_2, \dots) = (p_k + q_k)_{k \geq 0} \\ p \cdot q &= (p_0 q_0, p_1 q_1, p_2 q_2, \dots) = (p_k q_k)_{k \geq 0} \end{aligned}$$

Per ogni $r \in \mathbb{Q}$ indichiamo con (r) la successione costante $p = (p_k)_{k \geq 0}$ (dove per ogni $k \geq 0$ si ha $p_k = r$). È facilissimo verificare che la struttura algebrica $(\mathbb{Q}^{\mathbb{N}}, +, \cdot)$ è un *anello commutativo unitario*, con elemento neutro per la addizione la successione (0) , ed elemento neutro per la moltiplicazione la successione (1) .

Definizione 3.2.1 *Una successione $p \in \mathbb{Q}^{\mathbb{N}}$ si dice di Cauchy se per ogni $\epsilon \in \mathbb{Q}$, $\epsilon > 0$, esiste $N \geq 0$ tale che per ogni $h, k \geq N$ risulta $|p_h - p_k| \leq \epsilon$.*

Si noti che si tratta dell'ordinaria definizione di successione di Cauchy dell'analisi matematica (cf. [22], Def. 12.1, p. 91) riportata all'universo dei soli numeri razionali.

Indichiamo con $\mathcal{C}$ il sottoinsieme di $\mathbb{Q}^{\mathbb{N}}$ costituito dalle successioni di Cauchy; esso non è vuoto in quanto, ad esempio, le successioni costanti sono successioni di Cauchy.

Conviene subito osservare che

Teorema 3.2.2 *Una successione di Cauchy $(p_k)_{k \geq 0}$ è limitata, ossia esiste L tale che per ogni $k \geq 0$ si ha $|p_k| \leq L$.*

DIMOSTRAZIONE. (cf. anche [22], p. 92) Usando la definizione di successione di Cauchy con $\epsilon = 1$, tutti gli elementi da un certo indice N in poi distano da p_N al massimo di $\epsilon = 1$. Pertanto, poiché gli elementi p_k con $0 \leq k \leq N - 1$ sono in numero finito, basta prendere $L = \max\{|p_0|, \dots, |p_{N-1}|, |p_N| + 1\}$. $\square$

Definizione 3.2.3 *Una successione $p \in \mathbb{Q}^{\mathbb{N}}$ si dice convergente a 0 se per ogni $\epsilon \in \mathbb{Q}$, $\epsilon > 0$, esiste $N \geq 0$ tale che per ogni $k \geq N$ si ha $|p_k| \leq \epsilon$.*

Si noti che si tratta dell'ordinaria definizione di successione convergente a 0 dell'analisi matematica riportata all'universo dei soli numeri razionali.

Indichiamo con $\mathcal{C}_0$ il sottoinsieme di $\mathbb{Q}^{\mathbb{N}}$ costituito dalle successioni convergenti a 0, come ad esempio la successione costante uguale a 0, o la $(1/k)_{k \geq 1}$.

Osservazione 3.2.4 Una successione di Cauchy $p \in \mathbb{Q}^{\mathbb{N}}$ si dice *convergente*, con limite $r \in \mathbb{Q}$, se la successione $p - (r)$ converge a 0, essendo (r) la successione costante uguale a r ; ossia se per ogni $\epsilon \in \mathbb{Q}$, $\epsilon > 0$, esiste $N \geq 0$ tale che $k \geq N$ implica $|p_k - r| \leq \epsilon$.

Vedremo in seguito che l'insieme delle successioni di Cauchy convergenti allo stesso limite $r \in \mathbb{Q}$ “è” il numero r “pensato come reale”. Vi sono naturalmente successioni di Cauchy $p \in \mathbb{Q}^{\mathbb{N}}$ non convergenti ad alcun limite $r \in \mathbb{Q}$, come quella introdotta nell'Esempio 1.3.1.

3.2.1 Definizione dei numeri reali

Il teorema seguente genera $\mathbb{R}$ come anello quoziente (cf. [18], p. 221).

Teorema 3.2.5 *Si ha:*

- (i) $\mathcal{C}$ è un sottoanello commutativo ed unitario di $\mathbb{Q}^{\mathbb{N}}$;
- (ii) $\mathcal{C}_0$ è un ideale massimale di $\mathcal{C}$, e pertanto.³

il quoziente $\mathbb{R} = \mathcal{C}/\mathcal{C}_0$ è un corpo commutativo, detto il campo dei numeri reali.

DIMOSTRAZIONE. (i) $\mathcal{C}$ non è vuoto, ed è sufficiente verificare che la differenza ed il prodotto per componenti di successioni di Cauchy sono successioni di Cauchy. Per la differenza basta applicare il noto argomento del tipo “ $\epsilon/2$ ” usato nel teorema di analisi sul limite della somma di due successioni (cf. [22], Prop. 4.1. (j₁)).

Siano $p, q \in \mathcal{C}$. Dato $\epsilon \in \mathbb{Q}$, $\epsilon > 0$, consideriamo $\epsilon/2$: esistono indici $N, M \geq 0$ tali che per ogni $h, k \geq N$ risulti $|p_h - p_k| \leq \epsilon/2$, e per ogni $h, k \geq M$ risulti $|q_h - q_k| \leq \epsilon/2$. Per ogni $h, k \geq \max\{N, M\}$ si ha allora che $|(p_h - q_h) - (p_k - q_k)| = |p_h - p_k + q_k - q_h| \leq |p_h - p_k| + |q_h - q_k| \leq \epsilon/2 + \epsilon/2 = \epsilon$. Quindi $p - q \in \mathcal{C}$.

Per il prodotto basta applicare il noto argomento del tipo “ $\epsilon/(2L)$ ” usato nel teorema di analisi sul limite del prodotto di due successioni (cf. [22], Prop. 4.1. (j₃)).

Siano $p, q \in \mathcal{C}$. Entrambe le successioni sono limitate: esiste quindi un intervallo $[-L, L]$ cui appartengono tutti i termini sia di p che di q . Dato $\epsilon \in \mathbb{Q}$, $\epsilon > 0$, consideriamo $\epsilon/(2L)$: esistono indici $N, M \geq 0$ tali che per ogni $h, k \geq N$ risulti $|p_h - p_k| \leq \epsilon/(2L)$, e per ogni $h, k \geq M$ risulti $|q_h - q_k| \leq \epsilon/(2L)$. Per ogni $h, k \geq \max\{N, M\}$ si ha allora che $|p_h q_h - p_k q_k| = |p_h q_h - p_k q_h + p_k q_h - p_k q_k| = |p_h q_h - p_k q_h| + |p_k q_h - p_k q_k| = |p_h - p_k| |q_h| + |p_k| |q_h - q_k| \leq L |p_h - p_k| + L |p_k| |q_h - q_k| \leq \epsilon$. Quindi $pq \in \mathcal{C}$.

³ Cf. [18], Teorema 9.30, p. 222; [26], p. 150. Ricordiamo che il quoziente di un anello commutativo *non unitario* su un suo ideale massimale non è necessariamente un corpo: ad esempio l’ideale $4\mathbb{Z}$ dei multipli di 4 è massimale nell’anello $2\mathbb{Z}$ dei numeri pari, ma l’anello $2\mathbb{Z}/4\mathbb{Z} = \{[0], [2]\}$ *non* è un corpo poiché $[2] \cdot [2] = [4] = [0]$.

(ii) $\mathcal{C}_0$ è un ideale massimale⁴ di $\mathcal{C}$ perché la differenza di due successioni convergenti a 0 è una successione convergente a 0 (cf. [22], Prop. 4.1. (j₂)) e il prodotto di una successione convergente a 0 per una successione di Cauchy è una successione convergente a 0; infatti una successione di Cauchy è limitata ed il prodotto di una successione limitata per una successione che tende a 0 ha limite 0. Quindi $\mathcal{C}_0$ è un ideale, ed è un ideale proprio di $\mathcal{C}$ perché non tutte le successioni di Cauchy convergono a 0, ad esempio la successione $(3) = (3, 3, 3, \dots)$ è una successione di Cauchy che converge a 3.

La dimostrazione della massimalità di $\mathcal{C}_0$ richiede un po' d'astuzia. Sia J un ideale (non necessariamente proprio) di $\mathcal{C}$ che contiene propriamente l'ideale $\mathcal{C}_0$: cioè $\mathcal{C}_0 \subset J \subseteq \mathcal{C}$. Poiché $J \neq \mathcal{C}_0$, in J esiste almeno una successione di Cauchy che non converge a 0, sia $(w_k)_{k \geq 0}$. Naturalmente $(w_k)_{k \geq 0}$ può avere termini nulli, ma questi possono essere solo in numero finito, perché altrimenti $(w_k)_{k \geq 0}$ avrebbe una sottosuccessione costante uguale a 0, ed essendo di Cauchy, tutta la successione $(w_k)_{k \geq 0}$ dovrebbe convergere a 0 (ricordiamo che una successione di Cauchy con una sottosuccessione convergente converge allo stesso limite della sottosuccessione, cf. [22], p. 92). Quindi da un certo indice N in poi sarà $w_k \neq 0$.

Ora la successione $(w_k)_{k \geq N}$ ha tutti i termini non-nulli e quindi invertibili, e, come affermato in [22] (p. 96, nota 2) “si può allora dimostrare che la successione $(1/w_k)_{k \geq N}$ è di Cauchy”. Se $(w_k)_{k \geq N}$ fosse in J , essendo questo un ideale dovrebbe appartenervi anche il prodotto di $(w_k)_{k \geq N}$ con $(1/w_k)_{k \geq N}$. Ma tale prodotto non è altro che l'unità di $\mathcal{C}$; e pertanto J , come ogni ideale in cui sia presente l'unità di $\mathcal{C}$, coincide con $\mathcal{C}$ e quindi non ci sono ideali propri che contengono l'ideale $\mathcal{C}_0$ delle successioni convergenti a 0. Nulla però ci assicura che avendo buttato via i primi N elementi di $(w_k)_{k \geq 0}$ ci resti una successione ancora appartenente a J . Dobbiamo quindi procedere diversamente.

Introduciamo una successione $(u_k)_{k \geq 0}$ con termini *tutti nulli ad eccezione dei primi N* (quindi convergente a 0, quindi appartenente a $\mathcal{C}_0$, quindi appartenente a J). Ora a J appartiene anche la successione som-

⁴ Ricordiamo che un sottoinsieme I dell'anello commutativo unitario A è un *ideale* se: $(I, +)$ è un sottogruppo del gruppo abeliano $(A, +)$ e se per ogni x in I ed ogni a in A il prodotto $x \cdot a$ è sempre in I . Un ideale I è un ideale *proprio* se è un sottoinsieme proprio di A , cioè non coincide con A . Un ideale proprio è un *ideale massimale* se non è contenuto strettamente in nessun altro ideale proprio.

ma $v = u + w$. Possiamo scegliere i primi N termini non-nulli di u in modo che $v = u + w$ abbia tutti i termini non-nulli e quindi invertibili. Dimostriamo che la successione $(1/v_k)_{k \geq N}$ è di Cauchy.

Notiamo che v non è convergente a 0, perché allora lo sarebbe la differenza $w = v - u$. Quindi deve esistere un $\epsilon^* \in \mathbb{Q}$, $\epsilon^* > 0$ tale che per ogni $n \geq 0$ si trova un indice $k \geq n$ in corrispondenza del quale si ha

$$|v_k| > \epsilon^* \quad (3.1)$$

Utilizzando la definizione di successione di Cauchy con $\epsilon^*/2$, troviamo un indice $N \geq 0$ tale che per ogni $h \geq N$ ed ogni $k \geq N$ risulta

$$|v_h - v_k| \leq \epsilon^*/2 \quad (3.2)$$

Prendiamo un qualunque indice $j \geq N$: scegliendo $n = N$ nella (3.1) si trova un $k \geq N$ per cui $|v_k| > \epsilon^*$, e nel contempo, essendo $j, k \geq N$ possiamo usare la (3.2) e otteniamo $|v_j - v_k| \leq \epsilon^*/2$. Mettendo assieme, per un qualunque indice $j \geq N$ si ottiene:

$$|v_j| = |(v_j - v_k) + v_k| \geq |v_k| - |v_j - v_k| > \epsilon^* - \epsilon^*/2 = \epsilon^*/2$$

Ora per ogni coppia di indici $j', j'' \geq N$ si ha

$$\left| \frac{1}{v_{j'}} - \frac{1}{v_{j''}} \right| = \left| \frac{v_{j''} - v_{j'}}{v_{j'} v_{j''}} \right| \leq (2/\epsilon^*)^2 \cdot |v_{j'} - v_{j''}|$$

dal che segue subito che la successione $(1/v_k)_{k \geq N}$ è di Cauchy. Quindi la successione costante (1), che è il prodotto della successione $(v_k)_{k \geq N}$ appartenente all'ideale J e della successione $(1/v_k)_{k \geq N}$ che, essendo di Cauchy, appartiene all'anello $\mathcal{C}$, appartiene per definizione di ideale a J , e quindi $J = \mathcal{C}$. Abbiamo provato che $\mathcal{C}_0$ è un ideale massimale, il che completa la prova del teorema. $\square$

Osservazione 3.2.6 Nella costruzione di Méray–Cantor il problema algebrico più grosso è evidentemente l'invertibilità (rispetto alla moltiplicazione) di una successione di Cauchy. Infatti, se $x = (x_0, x_1, x_2, \dots)$ è una successione di Cauchy di razionali, per definire $1/x$, non basta eliminare eventuali componenti nulle, si pensi ad esempio alla successione di Cauchy

$$x = (1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{1}{5}, \dots, \frac{1}{k}, \dots)$$

che non ha componenti nulle, ma la cui inversa $1/x = (1, 2, 3, \dots, k, \dots)$ non è certo di Cauchy (perché non è limitata).

3.2.2 Immersione dei razionali

Per ogni numero razionale $r \in \mathbb{Q}$ possiamo costruire una successione costante (r) con ogni termine uguale ad r , e quindi considerare il numero reale

$$\hat{r} = (r) + \mathcal{C}_0$$

Viene così definita un'applicazione $\phi : r \in \mathbb{Q} \mapsto \hat{r} \in \mathbb{R}$ che è iniettiva perché la differenza $r - s$ fra due successioni costanti diverse non converge a 0; possiamo quindi *immergere* $\mathbb{Q}$ in $\mathbb{R}$ identificando $\mathbb{Q}$ con il sottoinsieme $\hat{\mathbb{Q}} = \phi(\mathbb{Q})$ di $\mathbb{R}$. In modo evidente ϕ è un omomorfismo, per cui $\phi : \mathbb{Q} \rightarrow \hat{\mathbb{Q}}$ è un isomorfismo.

3.2.3 La struttura d'ordine

Per introdurre l'ordinamento nel modello $\mathbb{R}$ di Méray–Cantor conviene usare il metodo dell'Osservazione 2.2.3.

Definizione 3.2.7 (Definizione dei numeri positivi) *Sia $a = (a_i) + \mathcal{C}_0$ un numero reale rappresentato dalla successione di Cauchy $(a_i)_{i \geq 0}$. Diciamo che a è positivo, in simboli,*

$$a > 0 \tag{3.3}$$

se esistono un $\epsilon^ \in \mathbb{Q}$ ed un indice $M \geq 1$ tali che per $n \geq M$*

$$a_n \geq \epsilon^* > 0 \tag{3.4}$$

Indichiamo con $\mathbb{R}^+$ il sottoinsieme degli elementi positivi.

La definizione necessita di una precisazione in quanto dobbiamo mostrare che essa non dipende dal rappresentante scelto. Supponiamo dati due rappresentanti di a , ad esempio sia $a = (a'_i) + \mathcal{C}_0 = (a_i) + \mathcal{C}_0$. Supponiamo che esistano un $\epsilon^* \in \mathbb{Q}$ ed un indice $M \geq 1$ tali che per $n \geq M$ risulti

$$a_n \geq \epsilon^* > 0 \tag{3.5}$$

La differenza $(a'_i) - (a_i)$ appartiene a $\mathcal{C}_0$, quindi, per definizione, esiste un indice $N \geq 1$ tale che per $n \geq N$ risulti

$$|a'_n - a_n| \leq \epsilon^*/2, \quad \text{cioè} \quad -\epsilon^*/2 \leq a'_n - a_n \leq \epsilon^*/2 \tag{3.6}$$

Allora, per indici $n \geq \max\{N, M\}$ risulta

$$a'_n = a_n + (a'_n - a_n) \geq \epsilon^* - \epsilon^*/2 = \epsilon^*/2 > 0 \quad (3.7)$$

□

Dimostriamo ora che:

Teorema 3.2.8 *Si ha:*

(P0) $0 \notin \mathbb{R}^+$

(P1) *Se $a, b \in \mathbb{R}^+$, allora $a + b, ab \in \mathbb{R}^+$.*

(P2) *Se $b \neq 0$, vale una ed una sola delle relazioni: $b \in \mathbb{R}^+$ oppure $-b \in \mathbb{R}^+$.*

DIMOSTRAZIONE. (P0) Il numero $0 \in \mathbb{R}$ è rappresentato da una qualunque successione di razionali convergente a 0, la quale non può verificare la (3.4).

(P1) Siano $a, b > 0$. Esistono $\epsilon_1^*, \epsilon_2^* \in \mathbb{Q}$ e due indici $M_1, M_2 \geq 1$ tali che per $n \geq M = \max\{M_1, M_2\}$ risulta

$$\begin{aligned} a_n &\geq \epsilon_1^* > 0 \\ b_n &\geq \epsilon_2^* > 0 \end{aligned}$$

da cui $a_n + b_n \geq \epsilon_1^* + \epsilon_2^* > 0$ e $a_n b_n \geq \epsilon_1^* \epsilon_2^* > 0$, e pertanto $a + b > 0$, $ab > 0$.

(P2) Ripeteremo un'argomentazione già usata nel Teorema 3.2.5. Sia $b \neq 0$. Pertanto una qualunque successione $(b_i)_{i \geq 0}$ che rappresenta b non converge a 0. Quindi deve esistere un $\epsilon^* \in \mathbb{Q}$, $\epsilon^* > 0$ tale che per ogni $n \geq 0$ si trova un indice $k \geq n$ in corrispondenza del quale si ha

$$|b_k| > \epsilon^* \quad (3.8)$$

Utilizzando la definizione di successione di Cauchy con $\epsilon^*/2$, troviamo un indice $N \geq 0$ tale che per ogni $h \geq N$ ed ogni $k \geq N$ risulta

$$|b_h - b_k| \leq \epsilon^*/2 \quad (3.9)$$

Prendiamo un qualunque indice $j \geq N$: scegliendo $n = N$ nella (3.8) si trova un $k \geq N$ per cui $|b_k| > \epsilon^*$, e nel contempo, essendo $j, k \geq N$

possiamo usare la (3.9) e otteniamo $|b_j - b_k| \leq \epsilon^*/2$. Ora k è un indice *fisso*, per cui vale una ed una sola delle relazioni:

$$b_k > \epsilon^* \quad \text{oppure} \quad -b_k > \epsilon^*$$

Nel caso $b_k > \epsilon^*$, mettendo assieme, per un qualunque indice $j \geq N$ si ottiene:

$$b_j = (b_j - b_k) + b_k \geq b_j - |b_j - b_k| > \epsilon^* - \epsilon^*/2 = \epsilon^*/2 > 0$$

cioè $b > 0$. Nel caso $-b_k > \epsilon^*$, per un qualunque indice $j \geq N$ si ottiene:

$$-b_j = (-b_j + b_k) - b_k \geq -b_j - |b_j - b_k| > \epsilon^* - \epsilon^*/2 = \epsilon^*/2 > 0$$

cioè $-b > 0$. □

Teorema 3.2.9 *Ponendo*

$$\begin{cases} a < b & \iff b - a \in \mathbb{R}^+ \\ a \leq b & \iff ((a < b) \text{ oppure } (a = b)) \end{cases}$$

si ottiene un ordine totale che fa di $\mathbb{R}$ un campo ordinato archimedeo, e che subordina su $\hat{\mathbb{Q}}$ l'ordine usuale.

DIMOSTRAZIONE. L'Osservazione 2.2.3 implica che $\mathbb{R}$ è un campo ordinato.

Verifichiamo la Proprietà di Archimede. Sia $x = (x_i) + \mathcal{C}_0 \in \mathbb{R}$. La successione $(x_i)_{i \geq 0}$ è di Cauchy, e quindi è limitata: in particolare esiste un intero $n > 0$ tale che risulti $x_i \leq n - 1$, ossia $n - x_i \geq 1$, per ogni $i \geq 0$. Dato che la successione costante $(n)_{i \geq 0}$ rappresenta il numero reale $n \cdot \hat{1}$, per la definizione della relazione d'ordine si ha $n \cdot \hat{1} > x$.

Dimostriamo infine che l'ordinamento subordina su $\hat{\mathbb{Q}}$ l'ordinamento usuale di $\mathbb{Q}$. Siano $r < s$ due razionali. I reali $\hat{r}$ e $\hat{s}$ sono rappresentati dalle successioni costanti $(r_n = r)_{n \geq 0}$ ed $(s_n = s)_{n \geq 0}$. Per verificare che $\hat{r} < \hat{s}$ basta prendere $\epsilon^* = s - r$; infatti in questo caso per ogni indice n si ha $r_n \leq s_n - (s - r)$. □

Osservazione 3.2.10 (Il valore assoluto) In accordo con la nota dell'Osservazione 2.3.7, il valore assoluto $|s|$ di un numero reale è definito ponendo $|s| = s$ se $s \geq 0$, $|s| = -s$ se $s \leq 0$. Osserviamo che se $s = (s_i) + \mathcal{C}_0$ con $(s_i)_{i \geq 0}$ una successione di Cauchy di numeri razionali, si ha $|s| = (|s_i|) + \mathcal{C}_0$.

3.2.4 Il teorema di completezza à la Méray–Cantor

Riprendiamo ora l'usuale definizione di successione di Cauchy in $\mathbb{R}$:

Definizione 3.2.11 Una successione $(s_i)_{i \geq 1}$ di numeri reali si dice di Cauchy se per ogni $\epsilon \in \mathbb{R}$, $\epsilon > 0$, esiste $N \geq 1$ tale che per ogni $n, m \geq N$ risulta

$$|s_n - s_m| \leq \epsilon \quad (3.10)$$

Si noti che la Definizione 3.2.11 contiene ben *quattro* quantificatori.⁵

Vediamo ora il teorema fondamentale della costruzione di Méray–Cantor: ricordiamo che

Definizione 3.2.12 Una successione $(s_i)_{i \geq 1}$ di numeri reali è convergente ad un numero reale σ se per ogni $\epsilon \in \mathbb{R}$, $\epsilon > 0$, esiste $N \geq 1$ tale che per ogni $n \geq N$ risulta

$$|s_n - \sigma| \leq \epsilon \quad (3.12)$$

Osserviamo che la Definizione 3.2.12 contiene *tre* quantificatori.⁶

Premettiamo al teorema fondamentale una importante proposizione, che afferma sostanzialmente che, se un numero reale σ è rappresentato da una successione di Cauchy $(v_m)_{m \geq 1}$ di numeri razionali, allora questa successione, “pensata” come successione di numeri reali, converge a σ .

⁵ In “formalese” la definizione si scrive

$$(\forall \epsilon \in \mathbb{R}, \epsilon > 0) (\exists N \geq 1) (\forall n \geq N) (\forall m \geq N) \quad (|s_n - s_m| \leq \epsilon) \quad (3.11)$$

⁶ In “formalese” la definizione si scrive

$$(\forall \epsilon \in \mathbb{R}, \epsilon > 0) (\exists N \geq 1) (\forall n \geq N) \quad (|s_n - \sigma| \leq \epsilon) \quad (3.13)$$

È stato osservato che nel linguaggio naturale si trovano tipicamente al massimo *due* quantificatori (come per esempio nel modo di dire: $\forall$ legge $\exists$ inganno), il che potrebbe essere uno degli ostacoli alla comprensione della definizione classica del limite delle funzioni reali, che contiene *tre* quantificatori ($\forall \epsilon \exists \delta \forall x \dots$). La definizione classica di successione di Cauchy di quantificatori ne ha addirittura quattro.

Proposizione 3.2.13 *Sia $\sigma = (v_m)_{m \geq 1} + \mathcal{C}_0$ un numero reale. Allora la successione di numeri reali $(\widehat{v_m})_{m \geq 1}$ converge a σ .*

DIMOSTRAZIONE. Sia dato $\epsilon \in \mathbb{R}$, $\epsilon > 0$. Dato che $\mathbb{R}$ è archimedeo, esiste un $\varepsilon \in \mathbb{Q}$, tale che $0 < \widehat{\varepsilon} < \epsilon$. Per essere $(v_m)_{m \geq 1}$ di Cauchy troviamo $N \geq 1$ tale che per $n, m \geq N$ sia $|v_n - v_m|_{\mathbb{Q}} \leq \varepsilon/2$. Fissando quindi $\delta^* = \varepsilon/2 > 0$ razionale, abbiamo che per ogni $n, m \geq N$

$$|v_n - v_m|_{\mathbb{Q}} \leq \varepsilon - \delta^*$$

Ora la successione costante $(v_n)_{m \geq 1}$ rappresenta $\widehat{v_n}$, mentre la $(v_m)_{m \geq 1}$ rappresenta σ per ipotesi; banalmente la successione costante $(\varepsilon)_{m \geq 1}$ rappresenta $\widehat{\varepsilon}$. Per la definizione di ordine questo significa che per ogni $n \geq N$ si ha

$$|\widehat{v_n} - \sigma|_{\mathbb{R}} < \widehat{\varepsilon} \quad (< \epsilon)$$

□

Dimostriamo ora, seguendo [24], il teorema fondamentale della costruzione di Méray–Cantor.

Teorema 3.2.14 *Ogni successione di Cauchy di numeri reali è convergente.*

DIMOSTRAZIONE. Indicheremo con $|\cdot|_{\mathbb{Q}}$ il valore assoluto in $\mathbb{Q}$ e con $|\cdot|_{\mathbb{R}}$ quello in $\mathbb{R}$. Sia $(s_i)_{i \geq 1}$ una successione di Cauchy di numeri *reali*, e per ogni i scegliamo una successione di Cauchy $(s_{im})_{m \geq 1}$ di numeri *razionali* che rappresenta s_i (scelta nella classe di equivalenza modulo $\mathcal{C}_0$ che è s_i). Sempre per ogni i , applicando la definizione di successione di Cauchy⁷ con $\epsilon = 1/(2i)$, troviamo un $N_i \geq 1$ tale che per ogni $m \geq N_i$ risulta:

$$|s_{im} - s_{i, N_i}|_{\mathbb{Q}} < \frac{1}{2i}$$

Definiamo la successione, di numeri razionali, “diagonale” $(v_i = s_{i, N_i})_{i \geq 1}$ e procediamo in quattro passi

⁷ Se una successione $(s_i)_{i \geq 1}$ è di Cauchy, allora, per definizione, per ogni $\epsilon > 0$, esiste $N \geq 1$ tale che per ogni $n, m \geq N$ risulta $|s_n - s_m| \leq \epsilon$, e quindi, *in particolare*, scegliendo $m = N$, per ogni $n \geq N$ risulta $|s_n - s_N| \leq \epsilon$.

- (a) Per ogni $i \geq 1$ si ha $|s_i - \widehat{v}_i|_{\mathbb{R}} < \widehat{1/i}$.

Il numero reale $|s_i - \widehat{v}_i|_{\mathbb{R}}$ a sinistra della disuguaglianza è rappresentato dalla successione di Cauchy razionale $(|s_{im} - s_{i,N_i}|_{\mathbb{Q}})_{m \geq 1}$. Infatti s_i è rappresentato dalla $(s_{im})_{m \geq 1}$; mentre il numero reale $\widehat{v}_i$ è rappresentato da una successione indicata in m costante, con tutti i termini uguali a $v_i = s_{i,N_i}$. Per $m \geq N_i$ si ha:

$$|s_{im} - s_{i,N_i}|_{\mathbb{Q}} < \frac{1}{2i} = \frac{1}{i} - \frac{1}{2i}$$

Esiste quindi $\epsilon^* = \frac{1}{2i} \in \mathbb{Q}$ ed esiste $M = N_i \geq 1$ tale che per ogni $m \geq M$

$$|s_{im} - s_{i,N_i}|_{\mathbb{Q}} \leq \frac{1}{i} - \epsilon^*$$

Per definizione di ordine, la tesi.

- (b) La successione “diagonale” di numeri razionali $(v_i)_{i \geq 1}$ è di Cauchy.

Fissiamo $\varepsilon \in \mathbb{Q}$, $\varepsilon > 0$. Usiamo la ipotesi che $(s_i)_{i \geq 1}$ è una successione di Cauchy di numeri reali, pertanto, esiste un indice J tale che, se $i \geq J$, si ha $|s_i - s_J|_{\mathbb{R}} < \widehat{\varepsilon}$. Poiché si può prendere $J > 2/\varepsilon$, per $i \geq J$ risulta

$$\begin{aligned} |\widehat{v}_i - \widehat{v}_J|_{\mathbb{R}} &= |\widehat{v}_i - \widehat{v}_J - s_i + s_i - s_J + s_J|_{\mathbb{R}} \\ &\leq |\widehat{v}_i - s_i|_{\mathbb{R}} + |\widehat{v}_J - s_J|_{\mathbb{R}} + |s_i - s_J|_{\mathbb{R}} \\ &\leq \widehat{1/i} + \widehat{1/J} + \widehat{\varepsilon} \\ &\leq \widehat{2\varepsilon} \end{aligned}$$

Quindi per $i \geq J$ si ha:

$$|v_i - v_J|_{\mathbb{Q}} \leq 2\varepsilon$$

- (c) Poniamo $\sigma = (v_m)_{m \geq 1} + \mathcal{C}_0$ (un tale σ è un numero reale).

- (d) Proviamo che $(s_i)_{i \geq 1}$ converge a σ .

Per ogni $i \geq 1$ risulta:

$$|s_i - \sigma|_{\mathbb{R}} \leq |s_i - \widehat{v}_i|_{\mathbb{R}} + |\widehat{v}_i - \sigma|_{\mathbb{R}}$$

Per la (a), si ha $|s_i - \widehat{v}_i|_{\mathbb{R}} < \widehat{1/i}$, e per la Proposizione 3.2.13 si ha $|\widehat{v}_i - \sigma|_{\mathbb{R}} \rightarrow 0$.

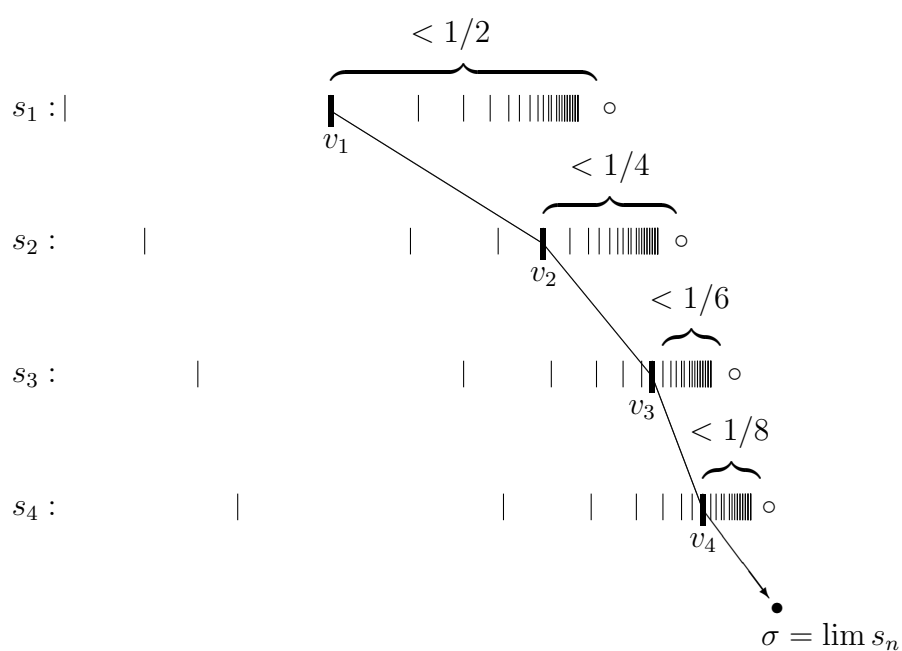


Figura 3.2: Illustrazione del teorema fondamentale della costruzione di Méray–Cantor. Le successioni $(s_{in})_{n \geq 1}$ sono per comodità disegnate come crescenti. I pallini vuoti $\circ$ indicano virtualmente i limiti in $\mathbb{R}$ di tali successioni (pensate con il cappello), ossia i reali $s_1, s_2, s_3, s_4, \dots$.

□

Dal Teorema 3.2.14 che assicura la convergenza in $\mathbb{R}$ delle successioni di Cauchy, segue il teorema di esistenza dell'estremo superiore.

Teorema 3.2.15 (Teorema di esistenza dell'estremo superiore)

Il modello $\mathbb{R}$ di Méray–Cantor è un campo ordinato completo.

DIMOSTRAZIONE. Sia X un sottoinsieme di $\mathbb{R}$ non-vuoto e superiormente limitato. Procediamo esattamente come nella dimostrazione del Teorema 2.7.2 di esistenza dell'estremo superiore a partire dalla Proprietà di Dedekind. Sia $a_0 \in X$ e sia b_0 una maggiorazione di X . Procediamo per bisezione.

- Consideriamo il punto di mezzo $t_0 = (a_0 + b_0)/2$: definiamo un intervallo

$$[a_1, b_1] = \begin{cases} [a_0, t_0] & \text{se } t_0 \text{ è una maggiorazione di } X \\ [t_0, b_0] & \text{se } t_0 \text{ non è una maggiorazione di } X \end{cases}$$

- Consideriamo il punto di mezzo $t_1 = (a_1 + b_1)/2$: definiamo un intervallo

$$[a_2, b_2] = \begin{cases} [a_1, t_1] & \text{se } t_1 \text{ è una maggiorazione di } X \\ [t_1, b_1] & \text{se } t_1 \text{ non è una maggiorazione di } X \end{cases}$$

- ...

L'algoritmo definisce due successioni $a = (a_n)$ e $b = (b_n)$ tali che $b_n - a_n = (b_0 - a_0)/2^n$. Per indici $m \geq n$ abbiamo per costruzione $a_n \leq a_m \leq b_m \leq b_n$, per cui le due successioni sono entrambe di Cauchy. Esistono pertanto un numero $\alpha \in \mathbb{R}$ al quale converge la $a = (a_n)$ ed un numero $\beta \in \mathbb{R}$ al quale converge la $b = (b_n)$.

Ripetendo l'argomentazione del teorema di analisi matematica del limite di una somma di successioni ([22], p. 63), $b_n - a_n \rightarrow \beta - \alpha$; si ha ([22], p. 57), $(b_0 - a_0)/2^n \rightarrow 0$; per il teorema di unicità del limite ([22], p. 61), si ha $\beta = \alpha$. Definiamo $\xi = \alpha = \beta$.

Proviamo che $\xi = \sup X$. Per ogni $x \in X$ e per ogni $n \geq 1$ si ha $x \leq b_n$ (essendo b_n una maggiorazione di X); per il teorema della permanenza del segno ([22], p. 65), si ha $x \leq \xi$. Inoltre, per ogni $\epsilon > 0$,

esistono a_n tali che $a_n > \xi - \epsilon$ (anzi, tutti gli a_n da un certo indice in poi verificano questa disuguaglianza); siccome a_n non è una maggiorazione di X , neppure $\xi - \epsilon$ può esserlo. $\square$

In conclusione, è stato dimostrato che

Teorema 3.2.16 *Il modello $\mathbb{R}$ di Méray–Cantor è un campo ordinato completo.*

Agli autori non sono note realizzazioni *in aula* di questo percorso didattico, ossia della definizione dei reali col modello di Méray–Cantor e con la dimostrazione del teorema di esistenza dell'estremo superiore come qui delineato. Si noti che sarebbe anche possibile un'introduzione assiomatica del tipo di Méray–Cantor, assumendo $\mathbb{R}$ un campo ordinato nel quale tutte le successioni di Cauchy convergono. Questo assioma può sembrare meno “credibile” della Proprietà di Dedekind o della stessa completezza (nel senso delle strutture ordinate), ma questo accade solo se la “credibilità” si confronta con l'intuizione geometrica, in particolare con quella della retta. Se invece ci si mette da un punto di vista più numerico, dal punto di vista ad esempio della rappresentazione decimale (che potrebbe esser conosciuta a livello “intuitivo” almeno tanto quanto la struttura della retta), la definizione di successione di Cauchy potrebbe essere espressa come segue:

Una successione $(x_n)_{n \geq 1}$ di numeri reali è di Cauchy se per ogni intero $k \geq 0$ esiste un indice $N \geq 1$ tale che per $n \geq N$ sia $|x_n - x_N| \leq 10^{-k}$.

Dal punto di vista della rappresentazione decimale “intuita” questo *non* significa che da N in poi tutti i numeri x_n hanno una rappresentazione decimale che coincide con una rappresentazione decimale di x_N fino alla $(k - 1)$ -esima cifra, neppure arrotondando. Ad esempio i numeri $x = 0.2399949$ e $y = x + 10^{-6}$ distano di 10^{-6} , i loro arrotondamenti⁸ a 5 cifre decimali sono rispettivamente 0.23999 e 0.24000, distano 10^{-5} , ed hanno uguali solo la prima cifra dopo la virgola. Tuttavia, l'essere di Cauchy significa per “molte” successioni la “stabilizzazione delle cifre”, per cui potrebbe essere abbastanza “credibile” un assioma di completezza che asserisca che tali successioni hanno limite. Per esempio qui sotto elenchiamo (arrotondate a 30 cifre) i primi 12 termini della successione

⁸ Nell'algoritmo di arrotondamento di R.

definita ricorsivamente da

$$x_1 = \pi, \quad \text{per } n \geq 2 \quad x_n = \frac{3(-8 - 4x_{n-1}^2 + x_{n-1}^4)}{4(-6 - 6x_{n-1} + x_{n-1}^3)}$$

che si dimostra essere di Cauchy:

```
|3.141592653589793238462643383279502884197
|6.082459565661631533130827527716219847448
4.|348269530332744472847964022181148970590
4.0|97677609244561593266957574470431869552
4.05|7770939144789404195590462647144839598
4.05680|9153791217030242625577650188153546
4.056808602850|650982868761743405627075899
4.0568086028504702646594148|62596292550298
4.056808602850470264659414843151820633579|
4.056808602850470264659414843151820633579|
4.056808602850470264659414843151820633579|
4.056808602850470264659414843151820633579|
```

La barra verticale | indica il punto dal quale le cifre si “stabilizzano”.

3.3 Commenti sulla definizione di successione di Cauchy

La nozione di successione di Cauchy è una sorta di *pons asinorum* nella costruzione dell’analisi. In questa sezione esaminiamo la definizione di successione di Cauchy, partendo da un “incidente” storico.

3.3.1 La definizione di Kline

Richiamiamo la definizione di successione di Cauchy come data da Cantor nell’articolo [11] già citato. Dobbiamo necessariamente esaminarla nella lingua originale:

Wenn ich von einer Zahlengrösse im weiteren Sinne rede, so geschieht es zunächst in dem Falle, dass eine durch ein Gesetz gegebene unendliche Reihe von rationalen Zahlen:

$$a_1, a_2, \dots, a_n, \dots \quad (1)$$

vorliegt, welche die Beschaffenheit hat, dass die Differenz $a_{n+m} - a_n$ mit wachsenden n unendlich klein wird, was auch die positive ganze Zahl m sei, oder mit anderen Worten, dass bei beliebig angenommen (positiven, rationalen) ε eine ganze Zahl n_1 vorhanden ist, so dass $|a_{n+m} - a_n| < \varepsilon$, wenn $n \geq n_1$ und wenn m eine beliebige positive ganze Zahl ist.⁹

Questa definizione originale di Cantor può essere scomposta a blocchi come segue:

... bei beliebig angenommen (positiven, rationalen) ε	eine ganze Zahl n_1 vorhanden ist,	so dass $ a_{n+m} - a_n < \varepsilon$,	wenn $n \geq n_1$ und wenn m eine beliebige positive ganze Zahl ist.
---	--------------------------------------	---	--

Si osservi la congiunzione “und” (ossia “e”), che lega in un unico blocco “wenn $n \geq n_1$ ” (“se $n \geq n_1$ ”) con “wenn m una beliebige positive ganze Zahl ist” (“se m è un intero positivo qualsiasi”). Si ha invece un’interpretazione errata se si legge *und* come congiunzione di tutto quanto viene prima con “wenn m una beliebige positive ganze Zahl ist”. In questo caso avremmo:

... bei beliebig angenommen (positiven, rationalen) ε eine ganze Zahl n_1 vorhanden ist, so dass $ a_{n+m} - a_n < \varepsilon$, wenn $n \geq n_1$	und wenn m eine beliebige positive ganze Zahl ist.
--	---

che può essere riscritto come:

⁹ *Trad. lib.*: Se dico che considero una grandezza numerica, nel senso generico del termine, ciò accade in primo luogo quando sia assegnata la legge di una successione infinita di numeri razionali $a_1, a_2, \dots, a_n, \dots$ che verifichi la condizione che la differenza $a_{n+m} - a_n$ sia infinitamente piccola al crescere di n quale che sia il numero positivo m , cioè, in altre parole, che, dato un qualsiasi ε (razionale positivo), esista un intero n_1 tale che sia $|a_{n+m} - a_n| < \varepsilon$ se $n \geq n_1$ e se m è un intero positivo qualsiasi.

$\lim_{n \rightarrow \infty} (a_{n+m} - a_n) = 0$	wenn m eine beliebige positive ganze Zahl ist.
---	--

il che è un errore di interpretazione in quanto i quantificatori, se sono posti *dopo* il predicato, come si usa talvolta nel linguaggio discorsivo,¹⁰ non possono essere disassociati e riassociati a piacere.¹¹ Questo errore interpretativo della definizione di successione di Cauchy si trova ad esempio in M. Kline [33], p. 984:

Cantor introduced a new class of numbers, ... starting with any sequence of rationals with obeys the condition that for any given ε , all the members except a finite number differ from each other by less than ε , or that

$$\lim_{n \rightarrow \infty} (a_{n+m} - a_n) = 0,$$

for arbitrary m . Such a sequence he calls a fundamental sequence. ... The next theorem is crucial. Cantor proves that if (b_ν) is any sequence of real numbers ... and if $\lim_{\nu \rightarrow \infty} (a_{\nu+\mu} - a_\nu) = 0$ for arbitrary μ , then there is a unique real number b , ... such that

$$\lim_{\nu \rightarrow \infty} b_\nu = a$$

Si noti che l'errore di traduzione è ripetuto¹² e si trova anche nella traduzione in italiano (cf. [34] vol. II, pp. 1148–1149) e anche in [39] (p. 275). In effetti la condizione

$$\lim_{n \rightarrow \infty} (a_{n+m} - a_n) = 0, \quad \text{for arbitrary } m \quad (3.14)$$

¹⁰ Si tratta di un uso assolutamente sconsigliabile nella didattica. La logica, o se si vuole, il linguaggio “formale” richiede naturalmente che i quantificatori vadano *prima* del predicato.

¹¹ Ricordiamo ad esempio che, dal punto di vista del linguaggio formale, la differenza fra la definizione di convergenza puntuale e quella di convergenza uniforme per una successione $f_n(x)$ di funzioni reali sta solo nello spostamento del quantificatore $\forall x$.

¹² Tuttavia in Kline la frase del linguaggio naturale “for any given ε , all the members except a finite number differ from each other by less than ε ” definisce correttamente le successioni di Cauchy: questa frase infatti afferma che per ogni dato $\varepsilon > 0$ esiste un indice N tale che tutti i termini con indici $\geq N$ (cioè tutti tranne un numero finito) differiscono due a due per meno di ε . L'errore interpretativo sta nelle parole che seguono “or that”.

ossia

$$(\forall m \geq 0) \quad \lim_{n \rightarrow \infty} |a_{n+m} - a_n| = 0 \quad (3.15)$$

non definisce le successioni di Cauchy, perché ad esempio è verificata dalla successione divergente $a_n = \log(n)$, in quanto per ogni $m \geq 0$ si ha

$$|a_{n+m} - a_n| = \log(n+m) - \log(n) = \log\left(1 + \frac{m}{n}\right) \rightarrow 0 \quad \text{per } n \rightarrow \infty$$

In effetti, se risolviamo rispetto ad n la disequazione:

$$\left| \log \frac{n+m}{n} \right| \leq \varepsilon$$

abbiamo:

$$n \geq N = N(\varepsilon, m) = \frac{m}{e^\varepsilon - 1}$$

e vediamo che N dipende sia da ε che da m . La dipendenza da m non può essere eliminata, poiché per ogni $\varepsilon > 0$ si ha $\sup_{m \geq 0} N(\varepsilon, m) = \infty$. Quindi, per ottenere le successioni di Cauchy, la formalizzazione di (3.14) dovrebbe essere implementata come:

$$\lim_{n \rightarrow \infty} |a_{n+m} - a_n| = 0 \quad \text{uniformemente in } m \geq 0 \quad (3.16)$$

che è in sostanza quella classica (3.2.11) a quattro quantificatori:

$$(\forall \epsilon \in \mathbb{R}, \epsilon > 0) (\exists N \geq 1) (\forall n \geq N) (\forall m \geq N) \quad (|s_n - s_m| \leq \epsilon) \quad (3.17)$$

3.3.2 La definizione di Bolzano

Il quarto quantificatore $\forall m$ nella (3.17), quello che, in quella posizione dove si trova, descrive l’“uniformemente in m ” della (3.16), può essere eliminato adottando la definizione seguente, talvolta attribuita a Bolzano, che usa solo *tre* quantificatori:

$$(\forall \epsilon \in \mathbb{R}, \epsilon > 0) (\exists N \geq 1) (\forall n \geq N) \quad (|s_n - s_N| \leq \epsilon) \quad (3.18)$$

Una definizione del tutto analoga si trova nel celebre *Trattato di algebra elementare: ad uso dei licei* di C. Arzelà, dove ([2], p. 343) si definisce una successione di Cauchy in $\mathbb{Q}$ come

...una serie di numeri razionali qualunque, ..., pei quali si possa nella serie medesima trovare un numero tale che la differenza tra esso e uno qualunque dei successivi sia in valore assoluto minore di un numero positivo scelto arbitrariamente piccolo ...

Discorsivamente, secondo Bolzano, una successione $(s_i)_{i \geq 1}$ di numeri reali è *di Cauchy* se per ogni $\epsilon > 0$ si trova un elemento “lontano” s_N nella successione tale che tutti gli elementi successivi distano da questo s_N al massimo ϵ .¹³ Proviamo la equivalenza fra (3.2.11) e la (3.18). Ponendo $m = N$ nella (3.2.11) si ottiene la (3.18). Quindi la (3.2.11) implica la (3.18). Per mostrare che la (3.18) implica la (3.2.11), basta scrivere la (3.18) con $\epsilon/2$ ed usare la disuguaglianza triangolare

$$|s_n - s_m| \leq |s_n - s_N| + |s_m - s_N|$$

Quindi le due definizioni sono equivalenti. $\square$

La definizione (3.18) di Bolzano ha due ulteriori proprietà di interesse didattico rispetto alla classica (3.2.11), oltre ad avere un quantificatore in meno:

- La definizione (3.2.12) di limite per una successione $(s_i)_{i \geq 1}$ di numeri reali, si ottiene semplicemente sostituendo un “numero fisso” σ all’elemento “lontano” s_N nella successione.
- È proprio nella forma di Bolzano che si usa la definizione di successione di Cauchy nel teorema fondamentale (3.2.14) della costruzione di Méray–Cantor.

3.3.3 Una definizione basata sulla finitezza

È prassi didattica usare l’avverbio “definitivamente” per indicare che una certa proprietà $P(n)$ che dipende da un indice intero n vale per tutti gli indici maggiori di un opportuno $\bar{n}$. Equivalentemente, la $P(n)$ vale definitivamente, se l’insieme degli indici per cui è falsa è un insieme finito. Per esempio, la definizione classica (3.2.12) di limite σ per una

¹³ È strano che, malgrado la diffusione in Italia del trattato dell’Arzelà, la definizione (3.18) a tre quantificatori sia stata poi sostituita nella pratica didattica dalla (3.2.11); parafrasando il celebre Rasoio di Occam, sarebbe da dire *Quantificatores non sunt multiplicandi praeter necessitatem*.

successione $(s_i)_{i \geq 1}$ di numeri reali equivale a

$$\forall \epsilon \in \mathbb{R}, \epsilon > 0, \text{ l'insieme } \{n \in \mathbb{N} \mid |s_n - s_m| > \epsilon\} \text{ è finito} \quad (3.19)$$

Ciò non può però essere esteso *verbatim* al caso delle successioni di Cauchy usando due indici: infatti, la condizione

$$\forall \epsilon \in \mathbb{R}, \epsilon > 0, \text{ l'insieme } \{(n, m) \in \mathbb{N} \times \mathbb{N} \mid |s_n - s_m| > \epsilon\} \text{ è finito} \quad (3.20)$$

definisce solo alcune successioni di Cauchy, precisamente quelle costanti. Infatti:

Proposizione 3.3.1 *La (3.20) implica che $(s_i)_{i \geq 1}$ è costante.*

DIMOSTRAZIONE. Supponiamo per assurdo che esistano indici i, j per i quali $s_i \neq s_j$, e sia $|s_i - s_j| = L > 0$. Scegliamo $\epsilon = L/4$. Per la (3.20), l'insieme delle coppie (n, m) di indici per i quali $|s_n - s_m| > L/4$ è un insieme finito, quindi contenuto in un “quadrato” Q di $\mathbb{N} \times \mathbb{N}$ di vertici sinistra–basso $(0, 0)$ e destra–alto $(\bar{n}, \bar{n})$, per un opportuno $\bar{n}$. Se fissiamo un indice $n > \bar{n}$, le coppie (i, n) ed (n, j) non appartengono a Q , e quindi si ha $|s_i - s_n| \leq L/4$ e $|s_n - s_j| \leq L/4$, quindi $|s_i - s_j| \leq |s_i - s_n| + |s_n - s_j| \leq L/4 + L/4 = L/2 < L$, assurdo. $\square$

Se consideriamo invece la condizione

$$\forall \epsilon \in \mathbb{R}, \epsilon > 0, \text{ l'insieme } \{\min\{n, m\} \in \mathbb{N} \mid |s_n - s_m| > \epsilon\} \text{ è finito} \quad (3.21)$$

si ha una definizione equivalente a quella di successione di Cauchy con addirittura un solo quantificatore; tutti gli altri sono stati “scaricati” sulla nozione di insieme. Si noti che, assegnato $\epsilon \in \mathbb{R}, \epsilon > 0$, per formare correttamente l'insieme $\{\min\{n, m\} \in \mathbb{N} \mid |s_n - s_m| > \epsilon\}$ dobbiamo per prima cosa considerare *per ogni* n e *per ogni* m le distanze $|s_n - s_m|$, e poi selezionare quegli interi n con $n \leq m$ per cui $|s_n - s_m| > \epsilon$. Il fatto che tali n siano in numero finito, significa che *per tutti gli indici* n *da un opportuno* $\bar{n}$ *in poi*, o è $n > m$, oppure, *quando* $n \leq m$ *risulta* $|s_n - s_m| \leq \epsilon$.

La definizione di successione di Cauchy nella forma (3.21) è esattamente quella utilizzata nel verificatore automatico di teoremi *Referee* (o *AetnaNova*) ideato principalmente da Eugenio Omodeo e Jacob T. Schwartz [47]. Il fatto che un software gestisca più facilmente la definizione (3.21) rispetto ad altre può suggerire interessanti spunti didattici per una sperimentazione in aula.

Capitolo 4

La costruzione di Dedekind

The devil is in the detail.^a

PROVERBIO POPOLARE

^a È nei dettagli che si annida il diavolo.



(a)



(b)

Figura 4.1: (a) R. Dedekind (1831–1916) e (b) Euclide (ca. 325–265 a.C.).

4.1 Introduzione

Questo capitolo presenta la costruzione dei numeri reali dovuta al matematico tedesco Richard Dedekind (1831–1916). Le laboriose verifiche for-

mali richieste da questa costruzione vengono presentate in dettaglio, mettendo in particolare evidenza i punti dove si usa la archimedeicità di $\mathbb{Q}$. Per il modello di $\mathbb{R}$ di Dedekind viene ovviamente provata la *completezza* nel senso dei corpi ordinati (esistenza dell'estremo superiore).

Dedekind sviluppò il suo modello nel 1858,¹ ma lo pubblicò solo nel 1872. I due contributi fondamentali di Dedekind in questa direzione, ossia *Was sind und was sollen die Zahlen* e *Stetigkeit und Irrationale Zahlen* sono stati raccolti e tradotti in italiano da Oscar Zariski [16], con l'aggiunta di interessantissime note storiche e critiche.

La quarta edizione in tedesco di *Stetigkeit und Irrationale Zahlen* (Friedr. Vieweg & Sohn, Braunschweig, 1912; 24 pp.) è resa disponibile nel web dalla Radboud University Nijmegen (Olanda) all'url:
www.ru.nl/w-en-s/gmfw/bronnen/dedekind2.html.

4.2 Le sezioni di Dedekind

La costruzione di Dedekind privilegia la relazione d'ordine (rispetto alla struttura algebrica), e culmina con il teorema di completezza.

4.2.1 Definizione dei numeri reali

Definizione 4.2.1 *Una sezione (Schnitt in tedesco, cut in inglese) di $\mathbb{Q}$ è una coppia (A, B) di sottoinsiemi di $\mathbb{Q}$ tali che*

(S0) $\{A, B\}$ è una partizione di $\mathbb{Q}$
(ossia $A \neq \emptyset$, $B \neq \emptyset$, $A \cap B = \emptyset$, $A \cup B = \mathbb{Q}$);

(S1) *Se $a \in A$ e $b \in B$ si ha $a < b$;*

(S2) *A non ha massimo.*

La definizione classica di sezione di Dedekind $x = (A, B)$ consiste soltanto nelle (S0) ed (S1), cosicché ci sono tre casi:

- A non ha massimo e B ha minimo r ;

¹ Dedekind ci indica anche la data precisa: il 24 novembre 1858.

- A non ha massimo e B non ha minimo;
- A ha massimo r e B non ha minimo;

Resta escluso il teorico quarto caso, ossia che A abbia massimo e B abbia minimo: infatti, se esistessero sia $a = \max A$ e $b = \min B$, per la densità di $\mathbb{Q}$ in sé, esisterebbe un razionale strettamente compreso fra a e b e quindi escluso da entrambe le classi, contro il fatto che $A \cup B = \mathbb{Q}$. In questa presentazione aggiungiamo (S2) alla definizione classica per evitare fastidiose triplicazioni nelle dimostrazioni.² Come si vedrà in seguito, la (S2), che esclude il terzo dei tre casi di cui sopra, lascia possibili *due casi*, che corrispondono ai *due casi* possibili per un numero reale: essere un razionale r , appunto il minimo di B , oppure essere irrazionale, ossia B senza minimo.³

Se $x = (A, B)$ è una sezione di $\mathbb{Q}$, per comodità di linguaggio chiameremo A la *classe inferiore* e B la *classe superiore* della sezione.

Proposizione 4.2.2 *Sia (A, B) una sezione di $\mathbb{Q}$. Allora:*

- A è inferiormente satura: per ogni $x \in \mathbb{Q}$, se $x < a \in A$ allora $x \in A$;
- B è superiormente satura: per ogni $x \in \mathbb{Q}$, se $x > b \in B$ allora $x \in B$;
- (A, B) è una coppia di classi contigue.

² Nel testo di analisi [22] che abbiamo preso come riferimento, viene usata (Definizione 3.1, p. 14) la definizione di sezione originale di Dedekind (S0)–(S1), senza assumere la (S2). Peraltro questo testo non include le dimostrazioni dettagliate della costruzione. Qui assumiamo (S2) nella definizione di sezione per evitare proprio le cosiddette “piccole difficoltà” di cui alla nota 11 in [22], p. 37.

³ Sostanzialmente agiamo come se nell’insieme delle sezioni di $\mathbb{Q}$ definite in base alla definizione originale di Dedekind (S0)–(S1) fosse introdotta un’equivalenza, definendo le sezioni (A, B) dove A non ha massimo e B non ha minimo equivalenti solo a se stesse, e definendo una sezione (A, B) dove A ha massimo r e B non ha minimo equivalente alla sezione $(A \setminus \{r\}, B \cup \{r\})$ dove A non ha massimo e B ha minimo r . Il ruolo della (S2) è in sostanza quello di scegliere esplicitamente un rappresentante canonico in ogni classe di equivalenza (scelta esplicita che è invece praticamente irrealizzabile nelle classi di equivalenza delle successioni di Cauchy considerate nel modello di Cantor–Méry).

DIMOSTRAZIONE. Mostriamo che A è inferiormente satura: sia $x \in \mathbb{Q}$ tale che $x < a \in A$. Per la (S0), se non fosse $x \in A$ dovrebbe essere $x \in B$, ma ciò è assurdo perché per la (S1) risulterebbe $a < x$. Analogamente B è superiormente satura: sia $x \in \mathbb{Q}$ tale che $x > b \in B$. Per la (S0), se non fosse $x \in B$ dovrebbe essere $x \in A$, una contraddizione con la (S1). Mostriamo infine che (A, B) è una coppia di classi contigue: sia $\epsilon \in \mathbb{Q}$, $\epsilon > 0$, e siano $a \in A$ e $b \in B$ arbitrari. Per la Proprietà di Archimede in $\mathbb{Q}$, esiste un intero $n \geq 1$ tale che $(b - a)/n < \epsilon$. Sia $x_k = a + k(b - a)/n$, con $k = 0, 1, \dots, n$, e sia

$$k^* = \max\{k \mid 0 \leq k \leq n, x_k \in A\}$$

Si ha $k^* < n$ in quanto $x_n = b \in B$. Poniamo $a^* = x_{k^*}$ e $b^* = x_{k^*+1}$, ovviamente $a^* \in A$ e $b^* \in B$ per la (S0) e la definizione di k^* . Si ha $b^* - a^* = (b - a)/n < \epsilon$. $\square$

Definizione 4.2.3 Una sezione $x = (A, B)$ di $\mathbb{Q}$ è detta numero reale. L'insieme di tutte le sezioni $x = (A, B)$ di $\mathbb{Q}$, ossia l'insieme di tutti i numeri reali, è indicato con $\mathbb{R}$.

4.2.2 Immersione dei razionali

Ad ogni numero razionale $r \in \mathbb{Q}$ possiamo associare una sezione $\hat{r} = (A, B)$ di $\mathbb{Q}$, definita da $A = \{q \in \mathbb{Q} \mid q < r\}$, $B = \{q \in \mathbb{Q} \mid q \geq r\}$. Viene così definita una applicazione $\phi : r \in \mathbb{Q} \mapsto \hat{r} \in \mathbb{R}$ iniettiva; possiamo quindi *immergere* $\mathbb{Q}$ in $\mathbb{R}$ identificando $\mathbb{Q}$ con il sottoinsieme $\phi(\mathbb{Q}) = \hat{\mathbb{Q}}$ di $\mathbb{R}$.

Il teorema seguente spiega il ruolo della classe superiore di una sezione.

Teorema 4.2.4 Sia (A, B) una sezione di $\mathbb{Q}$. Se B ha minimo r , allora $(A, B) = \hat{r}$.

DIMOSTRAZIONE. Sia r il minimo di B . Proviamo che $A = \{q \in \mathbb{Q} \mid q < r\}$. Se $q \in A$ allora per la (S1) deve essere $q < r \in B$; viceversa, se $q < r$, q non può stare in B (essendone r il minimo), quindi deve stare in A . Proviamo ora che $B = \{q \in \mathbb{Q} \mid q \geq r\}$. Se $q \in B$, si ha $q \geq r$ (per

definizione di minimo); viceversa, se $q \geq r$, q non può stare in A , per la (S1), e quindi deve stare in B . $\square$

Le sezioni di $\mathbb{Q}$ del tipo

$$\hat{r} = (\underbrace{\{q \in \mathbb{Q} \mid q < r\}}_A, \underbrace{\{q \in \mathbb{Q} \mid q \geq r\}}_B) \quad (4.1)$$

sono dette *sezioni di prima specie*, o numeri razionali. Le altre sono dette *sezioni di seconda specie*, o numeri irrazionali.

Osservazione 4.2.5 Intuitivamente, una sezione va vista come una spada che taglia $\mathbb{Q}$ con un fendente: la spada può “colpire” un razionale r , che poi viene buttato nella classe superiore, e sono i tagli di prima specie; oppure può tagliare $\mathbb{Q}$ senza colpire alcun punto, e sono i tagli di seconda specie.

Un esempio di sezione di $\mathbb{Q}$ di seconda specie

Proposizione 4.2.6 *Sia N un intero positivo non quadrato. Consideriamo*

$$A = \mathbb{Q}^- \cup \{0\} \cup \{a \in \mathbb{Q}^+ \mid a^2 < N\}, \quad B = \{b \in \mathbb{Q}^+ \mid b^2 > N\}$$

Allora (A, B) è una sezione di $\mathbb{Q}$ di seconda specie.

DIMOSTRAZIONE. Verifichiamo che (A, B) è una sezione. La (S0) vale perché l'equazione $x^2 = N$ nell'incognita x non ha soluzioni x razionali per come è stato preso N . Proviamo (S1): sicuramente gli elementi $a \leq 0$ di A sono minori degli elementi $b \in B$ che sono tutti positivi; sia allora $0 < a \in A$ e sia $b \in B$, si ha $a^2 < N < b^2$ da cui segue $a < b$ perché, se fosse $a \geq b$, avremmo $a^2 \geq ab \geq b^2$. Proviamo infine la (S2): supponiamo per assurdo che esista $a = \max(A)$. Risulta $a^2 < N$ e $a^2 + d = N$ con d razionale positivo. Si ha $d < 2a + 1$ perché, se fosse $d > 2a + 1$ si avrebbe $N = a^2 + d > a^2 + 2a + 1 = (a + 1)^2$ da cui $(a + 1) \in A$ che contraddice $a = \max(A)$; inoltre è $d \neq 2a + 1$ perché, se fosse $d = 2a + 1$, allora $N = a^2 + 2a + 1$ sarebbe un quadrato, contro l'ipotesi. Poiché $d < 2a + 1$ risulta

$$0 < \frac{d}{2a + 1} < 1$$

Posto

$$r = \frac{d}{2a+1}$$

si ha⁴ $2ar + r^2 < d$ da cui $a^2 + (2ar + r^2) < a^2 + d$, $(a+r)^2 < N$ e pertanto $a' = a + r \in A$ e ciò contraddice $a = \max(A)$.

Proviamo che B non ha minimo: supponiamo per assurdo che esista $b = \min(B)$. Risulta $b^2 > N$ e $b^2 - d = N$ con d razionale positivo. Sia

$$b' = b - \frac{d}{2b}$$

Si ha $0 < b' < b$ e

$$b'^2 = b^2 - d + \frac{d^2}{4b^2} > b^2 - d = N$$

da cui $b' \in B$ e ciò contraddice $b = \min(B)$.

Quindi (A, B) è di seconda specie. □

Per esempio, sia $N = 2$ e prendiamo $a = \frac{141}{100} = 1.41$, $b = \frac{142}{100} = 1.42$.
Si ha $a^2 = \frac{19881}{10000} = 1.9881 < 2$, e $b^2 = \frac{20164}{10000} = 2.0164 > 2$.

In base alle definizioni date si ha

$$a' = a + \frac{d}{2a+1} = \frac{53981}{38200} = 1.41312\dots$$

$$b' = b - \frac{d}{2b} = \frac{10041}{7100} = 1.41423\dots$$

$$a'^2 = \frac{2913948361}{1459240000} = 1.99689\dots < 2$$

$$b'^2 = \frac{100821681}{50410000} = 2.00003\dots > 2$$

Ricordiamo che $\sqrt{2} = 1.41421\dots$

4

$$\begin{aligned} d &< 2a+1 \\ d &< (2a+1)^2 - 2a(2a+1) \\ 2a(2a+1) + d &< (2a+1)^2 \\ 2ad(2a+1) + d^2 &< (2a+1)^2 d \\ 2a \frac{d}{2a+1} + \frac{d^2}{(2a+1)^2} &< d \\ 2ar + r^2 &< d \end{aligned}$$

4.2.3 La struttura d'ordine

Nella costruzione di Dedekind la relazione d'ordine è fondamentale. Introduciamola e vediamo le proprietà.

Definizione 4.2.7 (Definizione dell'ordine) Siano $x = (A, B)$, $y = (C, D)$ due numeri reali. Si definisce

$$x \leq y \iff A \subseteq C$$

Come di consueto, anche per i reali, $x < y$ equivale a $x \leq y$ con $x \neq y$.

Teorema 4.2.8 La relazione $\leq$ in $\mathbb{R}$ è un ordine totale, che subordina su $\hat{\mathbb{Q}}$ l'ordine usuale.

DIMOSTRAZIONE. La proprietà riflessiva è ovvia; anche l'antisimmetrica è ovvia, in quanto se $A = C$ si ha (per definizione di partizione) anche $B = D$; la proprietà transitiva è ovvia. Mostriamo che l'ordine è totale; siano $x = (A, B)$, $y = (C, D)$ due numeri reali qualsiasi. Se $A \subseteq C$ si ha $x \leq y$; altrimenti è $A \not\subseteq C$. Proviamo che allora $C \subset A$ (e quindi $y \leq x$). Esiste $q \in A$, $q \notin C$. Quindi $q \in D$ ed ogni elemento p di C è minore di q , e allora tale generico elemento p di C deve stare in A (se p stesse in B , dovrebbe essere $q < p$, mentre è $p > q$). Proviamo ora che se r, s sono razionali:

$$r < s \iff \hat{r} < \hat{s}$$

Per definizione delle sezioni di prima specie la disuguaglianza $\hat{r} < \hat{s}$ significa che

$$\{q \in \mathbb{Q} \mid q < r\} \subseteq \{q \in \mathbb{Q} \mid q < s\}$$

ovvero che per ogni $q \in \mathbb{Q}$ se $q < r$ allora $q < s$, e questo è ovviamente vero se $r < s$. $\square$

Teorema 4.2.9 (Densità dei razionali) ⁵ Siano $x = (A, B) < y =$

⁵ Nei campi ordinati questo teorema segue dalla Proprietà di Archimede: però, a questo punto della costruzione di Dedekind, non ha senso chiedersi se nel sistema $\mathbb{R}$ dei numeri reali di Dedekind vale o non vale la Proprietà di Archimede, in quanto non abbiamo ancora definito le operazioni algebriche; potremmo dimostrare facilmente che non ci sono numeri infiniti, ossia che per ogni $x = (A, B)$ reale esiste un naturale m tale che $\hat{m} > x$ (per esempio va bene un qualunque naturale della classe B), ma non possiamo formare i multipli $n\epsilon$ con un “quattrino” ϵ reale positivo arbitrario per superare lo “zecchino” x , né tantomeno scrivere la disequazione $n\epsilon > x$ nella forma $n > x/\epsilon$, per mancanza delle operazioni fra reali.

(C, D) due numeri reali. Esiste $q \in \mathbb{Q}$ tale che

$$x < \widehat{q} < y$$

DIMOSTRAZIONE. Proviamo che esiste un razionale q che soddisfa la condizione (i) $q \in C$, la condizione (ii) $q \notin A$, e, solo nel caso in cui $x = (A, B)$ sia una sezione di prima specie, anche la ulteriore condizione (iii) $q \neq \min B$. Infatti, dalla disuguaglianza $x < y$ segue che $A \subseteq C$, $A \neq C$, e pertanto si può trovare un razionale q che verifica la (i) e la (ii). Supponiamo che $x = (A, B)$ sia una sezione di prima specie, e che sia stato scelto proprio $q = \min B$: essendo che $q \in C$ e che C non ha massimo, esiste $q' \in C$, $q' > q$, e pertanto $q' \in B$, e $q' \notin A$. In questo caso quindi, q' verifica (i), (ii) e (iii).

Dalla (i), essendo la classe C inferiormente satura, si ha che ogni razionale minore di q appartiene a C , ovvero $\{r \in \mathbb{Q} \mid r < q\} \subseteq C$, e quindi, per definizione, $\widehat{q} \leq y$. È escluso che sia $\widehat{q} = y$, perché allora y sarebbe una sezione di prima specie con $q = \min D$, assurdo poiché $q \in C$.

Dalla (ii) si deduce che $q \in B$, e quindi ogni elemento di A è minore di q , ovvero $A \subseteq \{r \in \mathbb{Q} \mid r < q\}$, e quindi, per definizione, $x \leq \widehat{q}$. È escluso che sia $\widehat{q} = x$, perché allora x sarebbe una sezione di prima specie con $q = \min B$, assurdo per la condizione (iii). Si ha quindi $x < \widehat{q} < y$. $\square$

4.2.4 Il teorema fondamentale della costruzione

Definizione 4.2.10 Una sezione di $\mathbb{R}$ è una coppia $(\mathcal{A}, \mathcal{B})$ di sottoinsiemi di $\mathbb{R}$ tali che

(S0b) $\{\mathcal{A}, \mathcal{B}\}$ è una partizione di $\mathbb{R}$;

(S1b) Se $\alpha \in \mathcal{A}$ e $\beta \in \mathcal{B}$ si ha $\alpha < \beta$;

(S2b) $\mathcal{A}$ non ha massimo.

Come nel caso delle sezioni di $\mathbb{Q}$ diamo come esempio di sezioni di $\mathbb{R}$ le sezioni di prima specie: per ogni numero reale x , ponendo

$$(\mathcal{A}, \mathcal{B}) = (\underbrace{\{z \in \mathbb{R} \mid z < x\}}_{\mathcal{A}}, \underbrace{\{z \in \mathbb{R} \mid z \geq x\}}_{\mathcal{B}}) \quad (4.2)$$

otteniamo una sezione $(\mathcal{A}, \mathcal{B})$ di $\mathbb{R}$ di prima specie. Sono gli unici esempi possibili: vale infatti il seguente teorema fondamentale della costruzione di Dedekind:

Teorema 4.2.11 *Ogni sezione $(\mathcal{A}, \mathcal{B})$ di $\mathbb{R}$ è di prima specie.*

DIMOSTRAZIONE. Sia $(\mathcal{A}, \mathcal{B})$ una sezione di $\mathbb{R}$, e sia

$$A = \phi^{-1}(\mathcal{A} \cap \widehat{\mathbb{Q}}), \quad B = \phi^{-1}(\mathcal{B} \cap \widehat{\mathbb{Q}})$$

dove $\phi : \mathbb{Q} \rightarrow \mathbb{R}$ è l'applicazione che associa ad ogni numero razionale r la sezione $\widehat{r}$ di prima specie di $\mathbb{Q}$. Cioè gli elementi di A sono i razionali r tali che $\widehat{r} \in \mathcal{A}$, e analogamente, gli elementi di B sono i razionali r tali che $\widehat{r} \in \mathcal{B}$. Proviamo che (A, B) è una sezione di $\mathbb{Q}$.

Verifichiamo la (S0) della Definizione 4.2.1.

- A e B non sono vuote: infatti, presi ad esempio $\alpha', \alpha'' \in \mathcal{A}$, $\alpha' < \alpha''$, per la densità di $\widehat{\mathbb{Q}}$ troviamo $r \in \mathbb{Q}$ tale che $\alpha' < \widehat{r} < \alpha''$; si ha allora $r \in A$; analogamente si ragiona per B .
- $A \cap B = \emptyset$: infatti se esistesse $r \in A \cap B$, allora sarebbe $\widehat{r} \in \mathcal{A} \cap \mathcal{B}$ e ciò è assurdo.
- $A \cup B = \mathbb{Q}$: infatti se esistesse $r \in \mathbb{Q}$, $r \notin A \cup B$, allora $\widehat{r}$ sarebbe escluso sia da $\mathcal{A}$ che da $\mathcal{B}$ e ciò è assurdo.

Verifichiamo la (S1) della Definizione 4.2.1.

- Siano $p \in A$, $q \in B$; se per assurdo fosse $p \geq q$, avremmo

$$\{r \in \mathbb{Q} \mid r < q\} \subseteq \{r \in \mathbb{Q} \mid r < p\}$$

quindi $\widehat{p} \geq \widehat{q}$, assurdo in quanto essendo $\widehat{p} \in \mathcal{A}$, $\widehat{q} \in \mathcal{B}$, si violerebbe la (S1b) della Definizione 4.2.10.

Verifichiamo la (S2) della Definizione 4.2.1.

- Supponiamo sia $p = \max(A)$, poiché $\mathcal{A}$ non ha massimo, l'elemento $\widehat{p}$, non può essere il massimo di $\mathcal{A}$, e pertanto esiste $\alpha \in \mathcal{A}$ tale

che $\hat{p} < \alpha$. Per la densità di $\hat{\mathbb{Q}}$ in $\mathbb{R}$, esiste $\hat{q}$ tale che $\hat{p} < \hat{q} < \alpha$, ma allora $q \in A$ e $p < q$, e ciò contraddice $p = \max A$.

Rimane dunque provato che (A, B) è una sezione di $\mathbb{Q}$.

Poniamo $x = (A, B)$. Dimostriamo che x è il minimo di $\mathcal{B}$, per cui, ripetendo il ragionamento del Teorema 4.2.4, la sezione $(\mathcal{A}, \mathcal{B})$ di $\mathbb{R}$ risulta essere di prima specie.

- Proviamo che per ogni $\beta \in \mathcal{B}$ si ha $x \leq \beta$. Se β è il minimo di $\mathcal{B}$ il teorema è provato. Altrimenti esiste $\beta' \in \mathcal{B}$, con $\beta' < \beta$, ed esiste un $\hat{r} \in \mathcal{B}$, con $\beta' < \hat{r} < \beta$, ed $r \in B$. Basta dunque provare che $x \leq \hat{r}$. Le corrispondenti classi inferiori sono A per x e $C = \{q \in \mathbb{Q} \mid q < r\}$ per $\hat{r}$. Per la (S1) della Definizione 4.2.1, ogni elemento di A è minore di ogni elemento di B ed in particolare di r , cioè $A \subseteq C$.

- Proviamo infine che $x \in \mathcal{B}$. Per assurdo sia $x \in \mathcal{A}$. Poiché per la (S2b) della Definizione 4.2.10 la classe $\mathcal{A}$ non ha massimo, esiste $\alpha \in \mathcal{A}$, con $\alpha > x$, ed esiste un $\hat{r} \in \mathcal{A}$, con $\alpha > \hat{r} > x$, ed $r \in A$. Ciò è assurdo perché, se $r \in A$, deve essere $\hat{r} \leq x$. $\square$

Osservazione 4.2.12 Intuitivamente, la costruzione di Dedekind ha lo scopo di “suturare” gli “strappi” di $\mathbb{Q}$ corrispondenti alla “sezionabilità” in sezioni di seconda specie. Il teorema fondamentale della costruzione mostra che con la soluzione proposta da Dedekind, per $\mathbb{R}$ non ripresenta lo stesso problema: $\mathbb{R}$ non può essere sezionato in sezioni di seconda specie.

4.2.5 Il teorema di completezza à la Dedekind

Definizione 4.2.13 Sia $E \subseteq \mathbb{R}$: una *maggiorazione* di E è un numero $y \in \mathbb{R}$ tale che per ogni $x \in E$ sia $x \leq y$. $E \subseteq \mathbb{R}$ si dice *superiormente limitato* se esistono maggiorazioni di E .

Teorema 4.2.14 (Esistenza dell'estremo superiore) Sia $E \subseteq \mathbb{R}$ un insieme non vuoto e superiormente limitato. Allora l'insieme M delle maggiorazioni di E ha minimo m , che è l'estremo superiore $\sup(E)$ di E .

DIMOSTRAZIONE. Sia $\neg M$ il complementare $\mathbb{R} \setminus M$ di M . Consideriamo la coppia $(\neg M, M)$ e verifichiamo che è una sezione di $\mathbb{R}$ (per il Teorema

fondamentale 4.2.11 questo ci basta). Ora $M \neq \emptyset$ poiché E è superiormente limitato. Poiché E non è vuoto, esiste $x \in E$; se $E = \{x\}$ si ha $x = \sup(E)$ ed il teorema è provato. Altrimenti, esistono almeno due elementi in E , ed il minimo dei due non è una maggiorazione di E , e quindi $\neg M$ non è vuoto. Visto che, ovviamente, $\neg M \cap M = \emptyset$ e $\neg M \cup M = \mathbb{R}$, vale la (S0b).

Per provare la (S1b), siano $\alpha \in \neg M$, $\beta \in M$. Dalla definizione di M segue che α non è una maggiorazione di E , mentre β è una maggiorazione di E , e pertanto esiste $x \in E$ tale che $\beta \geq x > \alpha$.

Verifichiamo infine la (S2b), ossia che $\neg M$ non ha massimo. Se $\alpha \in \neg M$, tale α non è una maggiorazione di E . Esiste quindi $x \in E$ tale che $x > \alpha$, e, per il Teorema 4.2.9, esiste $\hat{r}$ tale che $x > \hat{r} > \alpha$. Quindi $\hat{r}$ non è una maggiorazione di E , e supera α . $\square$

4.2.6 La struttura algebrica

Per completare la costruzione di Dedekind dei numeri reali, dobbiamo definire le operazioni algebriche di addizione e moltiplicazione, provare che con tali operazioni si ottiene un *campo ordinato*. Le verifiche sono estremamente noiose. Per brevità per definire un numero reale x assegneremo solo la classe inferiore A della sezione $x = (A, B)$ (essendo la superiore B il complementare di A in $\mathbb{Q}$), e scriveremo $x := A$. Più precisamente:

Teorema 4.2.15 *Sia $A \subseteq \mathbb{Q}$ un sottoinsieme*

- (R1) *proprio,*
- (R2) *non-vuoto,*
- (R3) *privo di massimo,*
- (R4) *inferiormente saturo; ossia, per ogni $a' \in \mathbb{Q}$, da $a' < a \in A$ segue $a' \in A$.*

Allora, posto $B = \mathbb{Q} \setminus A$, la coppia (A, B) è una sezione. Viceversa, la classe inferiore A di una sezione verifica le quattro proprietà (R1)–(R4).

DIMOSTRAZIONE. Per la (R1) si ha $B \neq \emptyset$; essendo anche $A \neq \emptyset$ per la (R2), per definizione di B come complementare di A è verificata la (S0) della Definizione 4.2.1. Verifichiamo la (S1): sia $b \in B$; se esistesse $a \in A$ tale che $a \geq b$, escludendo che sia $a = b$ poiché $A \cap B = \emptyset$, avremmo

$a > b$, e quindi anche $b \in A$ (perché per la (R4) A è inferiormente saturo), assurdo. La (S2) infine coincide con (R3). Il viceversa è già noto dalla definizione e dalle ulteriori proprietà delle sezioni dimostrate in precedenza. $\square$

Il teorema precedente è legato al *metodo delle semirette di Russell* per la costruzione dei numeri reali, già citato nel Capitolo 2 (cf. [58], p. 220). Si tratta in sostanza del metodo di Dedekind limitato alla classe inferiore. Ricordiamo che, nel metodo di Russell, un sottoinsieme $A \subseteq \mathbb{Q}$ è detto una *semiretta di Russell* se è (R1') superiormente limitato, (R2) non-vuoto, (R3) privo di massimo, ed (R4) inferiormente saturo. Ora una semiretta di Russell, essendo superiormente limitata per la (R1') è un sottoinsieme proprio, ovvero (R1') implica (R1). Viceversa, un sottoinsieme $A \subseteq \mathbb{Q}$ che verifichi le quattro proprietà (R1)–(R4) è una semiretta di Russell: esiste infatti per la (R1) un razionale $b \notin A$, e tale b è una maggiorazione per A , perché b , se fosse superato da qualche elemento di A , dovrebbe anch'esso stare in A per la (R4).

Definiamo ora in $\mathbb{R}$ le operazioni di addizione e moltiplicazione e verifichiamo rispetto a queste due operazioni ed alla relazione d'ordine totale $\leq$ già introdotta, $\mathbb{R}$ è un campo ordinato. Per comodità del lettore seguiamo la numerazione di [22], pp. 9–10.

(A) **Ordine.** L'ordinamento totale $\leq$ è stato introdotto nella Sezione 4.2.7.

(B) **Addizione.** Se $\alpha := A$, $\beta := B$, definiamo

$$\alpha + \beta := \{a + b \mid a \in A, b \in B\}$$

La somma $\alpha + \beta$ è quindi definita semplicemente sommando le classi inferiori A di α e B di β . L'insieme $A + B = \{a + b \mid a \in A, b \in B\}$ è assunto come classe inferiore della somma. Tale classe $A + B$ non ha massimo. Se fosse $a + b = \max(A + B)$ (con $a \in A, b \in B$), non potendo essere né $a = \max A$, né $b = \max B$, troveremmo $a' > a$ in A e $b' > b$ in B tali che $a' + b' > a + b$ sia un elemento di $A + B$. Proviamo che $A + B$ è inferiormente saturo. Sia $z < a + b \in A + B$ (con $a \in A, b \in B$). Si ha $(a + b - z)/2 > 0$, $a - (a + b - z)/2 \in A$, $b - (a + b - z)/2 \in B$, sommando: $a - (a + b - z)/2 + b - (a + b - z)/2 = z \in A + B$. Quindi la somma $\alpha + \beta$ è definita correttamente, cioè è una sezione.

Abbiamo in precedenza identificato $\mathbb{Q}$ con il sottoinsieme $\widehat{\mathbb{Q}}$ di $\mathbb{R}$. Come primo passo per provare che $\phi : \mathbb{Q} \rightarrow \widehat{\mathbb{Q}}$ è un isomorfismo di campo, proviamo che per ogni $r, s \in \mathbb{Q}$

$$\widehat{r + s} = \widehat{r} + \widehat{s}$$

Ora $\widehat{r} := \{q \in \mathbb{Q} \mid q < r\}$, $\widehat{s} := \{q \in \mathbb{Q} \mid q < s\}$; inoltre la somma di queste due *semirette* razionali aperte è evidentemente la semiretta $\{q \in \mathbb{Q} \mid q < r + s\}$, ossia $\widehat{r + s}$.

- (B₁) **Proprietà commutativa dell'addizione.** Segue subito dall'analoga proprietà per sottoinsiemi: $A + B = B + A$.
- (B₂) **Proprietà associativa dell'addizione.** Segue subito dall'analoga proprietà per sottoinsiemi: $A + (B + C) = (A + B) + C$.
- (B₃) **Lo zero.** Sia

$$0 = \widehat{0} := \{q \in \mathbb{Q} \mid q < 0\}$$

Calcoliamo $\alpha + 0$. La classe inferiore A' della somma $\alpha + 0$ è la classe inferiore A di α più arbitrari numeri $q < 0$. Si ottengono numeri minori di numeri di A , quindi ancora in A , ossia $A' \subseteq A$. Viceversa, sia $a \in A$; ma A non ha massimo, quindi esiste $a' \in A$, $a' > a$; allora $a = a' + (a - a')$ ossia $a \in A'$, come somma di un elemento di A e di un $x = a - a' < 0$, ossia $A \subseteq A'$; quindi $A = A'$, ossia $\alpha + 0 = \alpha$.

- (B₄) **L'opposto.** Per ogni $\alpha = (A, B)$ definiamo:

$$-\alpha := \{q \in \mathbb{Q} \mid -q \in B, \text{ ma } -q \text{ non è il minimo di } B\}$$

ossia la classe inferiore A' del candidato opposto $-\alpha$ di α , è il simmetrico $-B$ della classe superiore B di α , privata eventualmente del suo minimo (perché altrimenti A' potrebbe avere massimo). Chiaramente A' è inferiormente satura, e quindi $-\alpha$ è definito correttamente, cioè è una sezione. Calcoliamo $\alpha + (-\alpha)$: la classe inferiore è $A + (-B) = \{a - b \mid a \in A, b \in B, b \neq \min B\}$, e sono tutti numeri < 0 . Per definizione della classe inferiore di 0, abbiamo che $\alpha + (-\alpha) \leq 0$. Viceversa dobbiamo ora provare che un qualunque $q < 0$ è uguale alla differenza fra un $a \in A$ ed un

$b \in B$. Poiché (A, B) è una coppia di classi contigue, esistono un $b \in B$ ed un $a \in A$ con differenza $b - a < (-q)/2$. Tenendo fisso b possiamo diminuire a fino a che $b - a = -q$. Si noti che in questa verifica relativa all'opposto è stata usata la proprietà di Archimede in $\mathbb{Q}$.

Per quanto riguarda la struttura moltiplicativa, conviene lavorare dapprima con i soli numeri reali positivi.⁶

(C) **Moltiplicazione di numeri positivi.** Indichiamo con $\mathbb{Q}^+$ l'insieme dei numeri razionali positivi e con $\mathbb{Q}^-$ l'insieme dei numeri razionali negativi. Se $\alpha := A$, $\beta := B$ sono entrambi > 0 definiamo

$$\alpha\beta := P = \mathbb{Q}^- \cup \{0\} \cup \{ab \mid a \in A \cap \mathbb{Q}^+, b \in B \cap \mathbb{Q}^+\}$$

Per provare che si tratta di una buona definizione, iniziamo col provare che tale classe P non ha massimo. Sia $p \in P$. Se è $p \leq 0$, tale p non è il massimo, essendovi in P numeri positivi. Se $p = ab > 0$, con $a > 0$ in A e $b > 0$ in B , basta maggiorare a in A e b in B per maggiorare il loro prodotto. Proviamo che P è inferiormente satura. Sia quindi $p \in P$ e sia $q < p$. Se $q \leq 0$ non ci sono problemi. Sia $q > 0$. Da $p > q > 0$ segue $p = ab$ con $a \in A$, $a > 0$, $b \in B$, $b > 0$. Scriviamo $q = ab(q/p)$; essendo $(q/p) < 1$ si ha $b(q/p) < b$ e quindi (per esser B inferiormente satura) $b(q/p) \in B$. Quindi q è il prodotto di $a \in A$, positivo, con $b(q/p) \in B$, positivo, e quindi sta in P .

Abbiamo in precedenza visto che per ogni $r, s \in \mathbb{Q}$ si ha $\widehat{r+s} = \widehat{r} + \widehat{s}$. Come secondo (ma non definitivo) passo per provare che $\phi : \mathbb{Q} \rightarrow \widehat{\mathbb{Q}}$ è un isomorfismo, proviamo che per ogni $r, s \in \mathbb{Q}$ positivi si ha

$$\widehat{rs} = \widehat{r}\widehat{s}$$

Ora $\widehat{r} := \{q \in \mathbb{Q} \mid q < r\}$, $\widehat{s} := \{q \in \mathbb{Q} \mid q < s\}$. Individuiamo la classe inferiore del prodotto $\widehat{r}\widehat{s}$: si tratta di

$$P = \mathbb{Q}^- \cup \{0\} \cup ([0, r[\cdot]0, s[)$$

⁶ Su questo aspetto della costruzione di Dedekind, vedi il commento nella sezione seguente.

dove gli intervalli aperti sono ovviamente intervalli razionali, e $]0, r[\cdot]0, s[$ è l'insieme di tutti i possibili prodotti di un elemento di $]0, r[$ ed un elemento di $]0, s[$. È ora evidente che $]0, r[\cdot]0, s[=]0, rs[$, per cui $P = \mathbb{Q}^- \cup \{0\} \cup]0, rs[$ è esattamente la classe inferiore di $\widehat{rs}$.

(C₁) **Proprietà commutativa della moltiplicazione.** Segue dal fatto che $ab = ba$ per a, b razionali.

(C₂) **Proprietà associativa della moltiplicazione.** Anche qui basta leggere la definizione del prodotto. Se α, β, γ sono reali positivi, per il prodotto $(\alpha\beta)\gamma$ avremo una classe inferiore con prodotti di razionali del tipo $(ab)c$, mentre per il prodotto $\alpha(\beta\gamma)$ avremo una classe inferiore con prodotti di razionali del tipo $a(bc)$, cioè le due classi sono uguali.

(C₃) **L'unità.** Sia

$$1 = \widehat{1} := \{q \in \mathbb{Q} \mid q < 1\}$$

Per $\alpha > 0$ calcoliamo $\alpha 1$. La classe inferiore P del prodotto $\alpha 1$ è

$$P = \mathbb{Q}^- \cup \{0\} \cup (\{a \mid a \in A \cap \mathbb{Q}^+\} \cdot]0, 1[)$$

dove A è la classe inferiore di α . Ora un $a \in A \cap \mathbb{Q}^+$ è un elemento positivo della classe inferiore A , e moltiplicato per un razionale positivo < 1 resta in A (che è inferiormente satura). Viceversa un qualunque elemento $a \in A \cap \mathbb{Q}^+$ è il prodotto di un $a' \in A \cap \mathbb{Q}^+$ per $a/a' < 1$, laddove sia $a' > a$ (A non ha massimo).

(C₄) **L'inverso.** Per ogni $\alpha = (A, B) > 0$ definiamo:

$$\alpha^{-1} := \mathbb{Q}^- \cup \{0\} \cup \{q \in \mathbb{Q} \mid 1/q \in B, \text{ ma } 1/q \text{ non è il minimo di } B\}$$

ossia la classe inferiore A' del candidato inverso α^{-1} di $\alpha > 0$ si ottiene invertendo gli elementi di B (positivi), dopo averne eventualmente rimosso il minimo (in modo che A' non abbia massimo), e saturando inferiormente con i razionali ≤ 0 .⁷ Chiaramente A' è

⁷ Conviene tracciare il grafico di $x \mapsto 1/x$ e vedere come si trasformano le semirette destre (aperte o chiuse).

inferiormente satura.⁸ Quindi α^{-1} è definito correttamente, cioè è una sezione. Calcoliamo $\alpha \alpha^{-1}$: la classe inferiore è

$$\alpha \alpha^{-1} := P = \mathbb{Q}^- \cup \{0\} \cup \{a/b \mid a \in A \cap \mathbb{Q}^+, b \in B \cap \mathbb{Q}^+, b \neq \min B\}$$

Dato che $a < b$, si ha $a/b < 1$ e quindi P è *contenuto* nella classe inferiore che definisce la unità, quindi $\alpha \alpha^{-1} \leq 1$. Resta da provare che P è *uguale* alla classe $\{q \in \mathbb{Q} \mid q < 1\}$. Sia dato un razionale q , $0 < q < 1$. Dobbiamo trovare $a \in A$ e $b \in B$ con rapporto $a/b = q$. Sappiamo che una sezione (A, B) di $\mathbb{Q}$ ha le classi contigue (grazie alla Proprietà di Archimede). Da questo segue che per ogni razionale q , $0 < q < 1$, esistono $a \in A$ e $b \in B$ tali che $q < a/b < 1$.⁹ Basta allora spingere b più in alto (B è superiormente satura) fino a $b' = b \cdot (a/b)/q$ per avere $a/b' = q$. Si noti che in questa verifica relativa all'inverso è *stata usata la proprietà di Archimede in $\mathbb{Q}$* .

(BC) **Proprietà distributiva per numeri positivi.** Anche qui basta applicare le definizioni. Siano $\alpha := A$, $\beta := B$, $\gamma := C$ reali positivi. La classe inferiore D' della somma $\alpha\gamma + \beta\gamma$ è costituita dalle somme di un elemento di

$$\{q \in \mathbb{Q} \mid q \leq 0\} \cup \{ac \mid a \in A, c \in C, a > 0, c > 0\}$$

e di uno di

$$\{q \in \mathbb{Q} \mid q \leq 0\} \cup \{bc \mid b \in B, c \in C, b > 0, c > 0\}$$

⁸ Sugeriamo di ricorrere ancora al grafico di $x \mapsto 1/x$

⁹ Supponiamo per assurdo che esista q^* , $0 < q^* < 1$, tale che per ogni $a \in A$ e ogni $b \in B$ riesca $a/b \leq q^* < 1$. Per ogni $n \in \mathbb{N}$, esistono elementi $a \in A$ e $b \in B$ tali che $b - a < 1/n$ e simultaneamente $a/b \leq q^* < 1$. Quindi possiamo scrivere

$$\frac{a}{b} = \frac{b + a - b}{b} = 1 - \frac{1}{nb} \leq q^* < 1$$

ossia

$$0 < 1 - q^* \leq \frac{1}{nb} < \frac{1}{na^*}$$

dove a^* è un qualsiasi elemento di A (quindi $a^* < b$). Si ottiene quindi una disuguaglianza valida per ogni $n \in \mathbb{N}$ (contenente solo n e le “costanti” q^* e a^* , senza più le variabili a, b) che viola evidentemente la Proprietà di Archimede.

ovvero dai razionali $ac + bc$ ($a \in A, b \in B, c \in C; a, b, c > 0$) e da tutti quelli più piccoli. La classe inferiore D'' del prodotto $(\alpha + \beta)\gamma$ è costituita dai razionali $(a + b)c$, con $a \in A, b \in B, c \in C, a + b > 0, c > 0$, e da tutti quelli più piccoli; tuttavia, quando facciamo variare $a \in A$ e $b \in B$ con la condizione che $a + b > 0$, basta prendere gli $a, b > 0$, in quanto con gli $a \leq 0$ e $b > 0$ oppure con gli $a > 0$ e $b \leq 0$ (fermo restando che $a + b > 0$) otteniamo somme minori di quelle che abbiamo con entrambi gli addendi positivi (tanto poi, dopo la moltiplicazione con $c > 0$, la classe viene saturata inferiormente).

(C bis) **Moltiplicazione di numeri a segno qualunque.** Il *valore assoluto* di α , indicato con $|\alpha|$, è

$$|\alpha| = \begin{cases} \alpha & \text{se } \alpha \geq 0 \\ -\alpha & \text{se } \alpha < 0 \end{cases}$$

In generale definiamo la moltiplicazione di numeri di segno qualunque:

$$\alpha\beta = \begin{cases} 0 & \text{se } \alpha = 0 \text{ oppure } \beta = 0 \\ |\alpha| |\beta| & \text{se } \alpha > 0, \beta > 0 \text{ oppure } \alpha < 0, \beta < 0 \\ -(|\alpha| |\beta|) & \text{se } \alpha > 0, \beta < 0 \text{ oppure } \alpha < 0, \beta > 0 \end{cases}$$

e, se $\alpha < 0$, l'inverso:

$$\alpha^{-1} = -(|\alpha|)^{-1}$$

La verifica delle proprietà commutativa e associativa per la moltiplicazione di numeri di segno qualunque, la verifica che per ogni $\alpha \neq 0$ si ha $\alpha\alpha^{-1} = 1$, e la verifica della proprietà distributiva per numeri di segno qualunque, sono ora di pura routine. La verifica che $\widehat{r}\widehat{s} = \widehat{r}\widehat{s}$ vale per razionali r, s di segno arbitrario discende subito dal fatto che in $\mathbb{Q}$ vale per la moltiplicazione la “regola dei segni”. Si è quindi fin qui provato che $\mathbb{R}$ è un campo, estensione di $\mathbb{Q}$. Resta da provare che:

(AB) **Assioma 1 di campo ordinato.** Proviamo che per ogni $\alpha := A, \beta := B, \gamma := C$ in $\mathbb{R}$ si ha che $\alpha \leq \beta \Rightarrow \alpha + \gamma \leq \beta + \gamma$. Per definizione dell'ordine si ha $A \subseteq B$. Evidentemente $A + C \subseteq B + C$.

(AC) **Assioma 2 di campo ordinato** Proviamo che per ogni $\alpha := A$, $\beta := B$, $\gamma := C$ in $\mathbb{R}$, se $\gamma > 0$ si ha che $\alpha \leq \beta \Rightarrow \alpha\gamma \leq \beta\gamma$. Se $A \subseteq B$, l'inclusione viene conservata moltiplicando per $c \in C$, $c > 0$. $\square$

In conclusione, è stato dimostrato che

Teorema 4.2.16 *Il modello $\mathbb{R}$ di Dedekind è un campo ordinato completo.*

4.3 Un percorso alternativo

Abbiamo visto che il Teorema 4.2.11 implica la completezza del modello $\mathbb{R}$ di Dedekind (Teorema 4.2.14). Questo è il percorso didattico “classico” (cf. ad esempio [19], pp. 260, 277). Infatti, come discusso nella Sezione 2.8, il Teorema 4.2.11 è la realizzazione concreta nel modello di Dedekind dell'*Axiom der Vollständigkeit*, e quindi il percorso didattico fin qui presentato corrisponde all'implicazione *Axiom der Vollständigkeit* $\Rightarrow$ *completezza* del Teorema 2.6.12. Il fatto che nel Teorema 2.6.12 valga anche l'implicazione opposta, cioè *completezza* $\Rightarrow$ *Axiom der Vollständigkeit*, suggerisce che nel modello $\mathbb{R}$ di Dedekind si possa dimostrare *prima* la completezza e *poi* far discendere da questa il teorema fondamentale della costruzione.

Ora non solo questo è possibile, ma, come vedremo, questo percorso alternativo, in talune situazioni, può apparire didatticamente più interessante di quello “classico”. Infatti:

- La dimostrazione diretta della completezza del modello $\mathbb{R}$ di Dedekind è molto naturale, è basata direttamente sulla definizione della relazione d'ordine, ed è il seguito logico della discussione euristica presentata nella Sezione 1.3.
- Questo percorso alternativo è adatto nelle situazioni in cui si sia interessati solo all'aspetto “tecnico” della costruzione, cioè se interessa solo provare la completezza del modello $\mathbb{R}$ di Dedekind rinunciando agli aspetti storici e “strutturali” legati all'*Axiom der Vollständigkeit*; in particolare questo percorso è quello adatto per l'introduzione di $\mathbb{R}$ basata sugli allineamenti di cifre, che (come

si vedrà nel Capitolo 6) è nient'altro che una “semplificazione” algoritmica della costruzione di Dedekind.

Vediamo quindi questo percorso alternativo.

Teorema 4.3.1 (Esistenza dell'estremo superiore) *Sia $E \subseteq \mathbb{R}$ non vuoto e superiormente limitato. Allora esiste l'estremo superiore $\sup(E)$ di E .*

DIMOSTRAZIONE. Per ogni $x \in E$ indichiamo con A_x la classe inferiore di x , e poniamo per definizione $A = \cup_{x \in E} A_x$. Dal Teorema 4.2.15 segue che $\xi = (A, \mathbb{R} \setminus A)$ è una sezione di $\mathbb{Q}$. Si vede infatti facilmente che $A \subseteq \mathbb{Q}$ è un sottoinsieme proprio di $\mathbb{Q}$ (data una maggiorazione m di E , i razionali della classe superiore di m non possono appartenere ad A), che non è vuoto (le classi A_x non sono vuote), che non ha massimo (dato $a \in A$, esiste una classe A_x alla quale a appartiene; dato che A_x non ha massimo, esiste $a' \in A_x \subseteq A$ tale che $a' > a$) e che è inferiormente saturo (tutte le classi A_x lo sono). Proviamo ora che $\xi = \sup E$. Per ogni $x \in E$ si ha $A_x \subseteq A$, ossia, per la definizione della relazione d'ordine, si ha $x \leq \xi$. Dunque ξ è una maggiorazione di E in $\mathbb{R}$. Per provare che ξ è la *minima* maggiorazione di E , basta ora mostrare che nessun reale $\eta < \xi$ può essere una maggiorazione di E . Sia quindi dato $\eta < \xi$. Per il Teorema 4.2.9 esistono *due* razionali $r < q$ tale che $\eta < \hat{r} < \hat{q} < \xi$. Evidentemente r appartiene alla classe inferiore di $\hat{q}$ e quindi, per la definizione della relazione d'ordine, r appartiene alla classe inferiore di ξ , ossia $r \in A$. Per la definizione di A , esiste $x \in E$ tale che $r \in A_x$, e pertanto si ha $\eta < \hat{r} < x$, il che esclude che η possa essere una maggiorazione di E . $\square$

Possiamo ora ridimostrare molto facilmente il teorema fondamentale della costruzione di Dedekind, usando direttamente le definizioni di estremo superiore e di sezione.

Teorema 4.3.2 *Ogni sezione $(\mathcal{A}, \mathcal{B})$ di $\mathbb{R}$ è di prima specie.*

DIMOSTRAZIONE. Sia $\xi = \sup \mathcal{A}$. Per la (S2b) della Definizione 4.2.10 delle sezioni di $\mathbb{R}$, la classe $\mathcal{A}$ non ha massimo, e pertanto $\xi \notin \mathcal{A}$ (altrimenti sarebbe $\xi = \max \mathcal{A}$). Per la (S0b) deve quindi essere $\xi \in \mathcal{B}$. Per la (S1b) ogni elemento di $\mathcal{B}$ è una maggiorazione di $\mathcal{A}$, e dato che, per definizione, ξ è la *minima* maggiorazione di $\mathcal{A}$, per ogni $\beta \in \mathcal{B}$ deve essere $\xi \leq \beta$. Abbiamo provato che $\xi = \min \mathcal{B}$ e pertanto la sezione $(\mathcal{A}, \mathcal{B})$ di $\mathbb{R}$ è di prima specie. $\square$

4.4 Commenti

Non sfugge certo al lettore che la costruzione di Dedekind è alquanto artificiosa dal punto di vista algebrico. Essa procede snella e abbastanza veloce fino all'introduzione dell'ordine, poi, con le operazioni algebriche, ed in particolare con la moltiplicazione, ha una vistosa caduta di eleganza. Questi "inestetismi" potrebbero essere collegati al fatto che la costruzione è essenzialmente ed implicitamente geometrica, tanto che la parte algebrica scorre abbastanza liscia per i numeri reali positivi, che sono *misure* di segmenti. Infatti, la definizione di sezione di Dedekind si trova "identica", a parte ovviamente la terminologia, nella famosa definizione di *rapporto fra grandezze* contenuta implicitamente nella Definizione 5 del Libro V degli Elementi di Euclide:

Quattro grandezze si dicono nello stesso rapporto, la prima alla seconda e la terza alla quarta, quando, se si prende un qualunque equimultiplo comune alla prima e alla terza, e un qualunque equimultiplo comune alla seconda e alla quarta, i primi equimultipli entrambi eccedono, o entrambi sono uguali, o entrambi minori dei secondi equimultipli presi rispettivamente nello stesso ordine.

In altre parole, date quattro grandezze A, B, C, D , si dice che $A : B = C : D$ (uguaglianza di rapporti, o proporzione) se per ogni coppia (n, m) di interi positivi si ha che

$$\begin{cases} mA > nB \implies mC > nD \\ mA = nB \implies mC = nD \\ mA < nB \implies mC < nD \end{cases} \quad (4.3)$$

Mostriamo come questa definizione equivalga alla considerazione di una sezione di Dedekind. Siano date due grandezze omogenee A, B tali che un multiplo dell'una qualunque delle due superi l'altra.¹⁰ Ogni coppia (n, m) di interi positivi appartiene ad uno ed uno solo dei seguenti tre sottoinsiemi: $H = \{(n, m) \mid mA > nB\}$, $J = \{(n, m) \mid mA = nB\}$, $K = \{(n, m) \mid mA < nB\}$, che restano i medesimi se, in luogo di A, B , si considerano grandezze C, D tali che $A : B = C : D$. Si noti che per ipotesi $H \neq \emptyset$, $K \neq \emptyset$, mentre può essere $J = \emptyset$. Identificando le coppie (n, m) con i rapporti $\frac{n}{m}$, si ottiene una partizione di $\mathbb{Q}^+$ in due classi

$$R = \left\{ \frac{n}{m} \in \mathbb{Q}^+ \mid nB < mA \right\}, \quad L = \left\{ \frac{n}{m} \in \mathbb{Q}^+ \mid nB \geq mA \right\}$$

¹⁰ Altrimenti la definizione euclidea di proporzione non ha senso.

tali che, se $a \in L$ e $b \in R$, risulta $a < b$. Quindi (L, R) è una sezione di Dedekind di $\mathbb{Q}^+$, ed identifica, secondo la attuale terminologia, un numero reale positivo, che altro non è, secondo la terminologia di Euclide, che il rapporto $A : B$ fra le grandezze date.

L'approccio di Dedekind non avrebbe quindi alcuna "originalità" scientifica, ma soltanto "filosofica", visto che la sua novità consiste nell'eliminazione di ogni considerazione di tipo geometrico-intuitivo nella costruzione dei numeri reali.

Capitolo 5

La completezza

Es ist nicht schwer, zu komponieren. Aber es ist fabelhaft schwer, die überflüssigen Noten unter den Tisch fallen zu lassen.^a

JOHANNES BRAHMS

^a Non è difficile comporre, ma è incredibilmente difficile far tacere le note superflue.

5.1 Introduzione

In questo capitolo mostriamo come il Teorema di Cauchy (ogni successione reale di Cauchy converge), che è il “coronamento” della costruzione di Méray–Cantor, ed il Teorema di Bolzano–Weierstrass (essenzialmente la compattezza sequenziale dei chiusi limitati), che a sua volta è il “coronamento” della costruzione di Dedekind, siano dimostrabili l’uno dall’altro, aprendo la possibilità di percorsi didattici diversi per iniziare lo studio dell’analisi matematica, a seconda della definizione di $\mathbb{R}$ adottata.

5.2 Il circolo della completezza

Sappiamo che il modello di Méray–Cantor fornisce la “completezza” di $\mathbb{R}$ nel senso tecnico della teoria degli spazi metrici: ogni successione di Cau-

chy è convergente, mentre la costruzione di Dedekind fornisce la “completezza” di $\mathbb{R}$ nel senso tecnico della teoria degli insiemi totalmente ordinati: ogni sottoinsieme non-vuoto e superiormente limitato ha estremo superiore.

Nel caso di un approccio assiomatico, è indifferente dare l’assioma di completezza in uno o nell’altro modo: infatti vediamo qui sotto un “circolo” di teoremi che mostra che i due risultati sono equivalenti.

Nel caso di approccio assiomatico, la prassi didattica più comune prende il Teorema 5.2.1 come assioma di completezza, e segue il circolo fino al Teorema 5.2.5. Sarebbe possibile prendere il Teorema 5.2.5 come assioma di completezza e poi percorrere il circolo fino all’importante Teorema 5.2.4. Come già segnalato, agli autori non sono note realizzazioni *in aula* di questo percorso didattico, forse perché si ritiene più intuitiva, o comunque meno complessa da enunciare, la esistenza dell’estremo superiore di quanto lo sia la convergenza delle successioni di Cauchy.

Prendiamo quindi come punto di partenza il Teorema 4.2.14 di completezza *à la* Dedekind, e sia quindi $\mathbb{R}$ il modello dei reali di Dedekind.

Teorema 5.2.1 (Esistenza dell’estremo superiore) *Sia X un sottoinsieme di $\mathbb{R}$ non-vuoto e superiormente limitato; allora esiste $\xi = \sup(X)$.*

L’esistenza dell’estremo superiore implica il teorema di Cantor:

Teorema 5.2.2 (Cantor) *Sia $(I_n)_{n \in \mathbb{N}}$ una successione di intervalli chiusi, decrescente per inclusione; allora $\bigcap_{n \in \mathbb{N}} I_n \neq \emptyset$*

DIMOSTRAZIONE. Posto $I_n = [a_n, b_n]$, l’insieme $\{a_n \mid n \in \mathbb{N}\}$ è non vuoto e superiormente limitato (da un qualunque elemento della successione $(b_n)_{n \in \mathbb{N}}$). Esiste quindi $\xi = \sup\{a_n \mid n \in \mathbb{N}\}$. Proviamo che $\xi \in \bigcap_{n \in \mathbb{N}} I_n$. Per la I PCES, per ogni $n \in \mathbb{N}$ si ha $a_n \leq \xi$. Dobbiamo provare che per ogni $n \in \mathbb{N}$ si ha $\xi \leq b_n$. Se ci fosse un n per cui $\xi > b_n$, per la II PCES troveremmo un m per cui $b_n \leq a_m$, e quindi $a_n < b_n \leq a_m < b_m$, un assurdo. $\square$

Il teorema di Cantor implica il teorema di Bolzano–Weierstrass:

Teorema 5.2.3 (Bolzano–Weierstrass) *Sia X un sottoinsieme di $\mathbb{R}$ infinito e limitato; allora X ha almeno un punto di accumulazione.*¹

DIMOSTRAZIONE. Si procede per bisezione di un intervallo $I_0 = [a_0, b_0]$ che contiene X , scegliendo al passo n l'intervallo $I_n = [a_n, b_n]$ che contiene infiniti punti di X . Sia $\xi \in \bigcap_{n \in \mathbb{N}} I_n$ (teorema di Cantor). Prendiamo un intorno $U = (\xi - \epsilon, \xi + \epsilon)$ di ξ . Un intervallo $[a_n, b_n]$ di ampiezza $(b_0 - a_0)/2^n < \epsilon$ dovendo contenere ξ deve essere interamente contenuto in U . $\square$

Il teorema di Bolzano–Weierstrass implica
che ogni successione limitata in $\mathbb{R}$ ha una sottosuccessione
convergente:

Teorema 5.2.4 *Sia $(x_n)_{n \in \mathbb{N}}$ una successione limitata di numeri reali; allora $(x_n)_{n \in \mathbb{N}}$ ha almeno una sottosuccessione convergente.*

DIMOSTRAZIONE. Se $X = \{x_n\}_{n \in \mathbb{N}}$ è un insieme finito, estraiamo una successione costante, se X è un insieme infinito e a è un suo punto di accumulazione, scegliamo un punto x_{n_k} nell'intorno $(a - 1/k, a + 1/k)$, ed otteniamo la successione $(x_{n_k})_{k \in \mathbb{N}}$ che converge ad a . $\square$

Il Teorema 5.2.4 implica che $\mathbb{R}$ è uno spazio metrico completo:²

Teorema 5.2.5 (Completezza metrica di $\mathbb{R}$) *Sia $(x_n)_{n \in \mathbb{N}}$ una successione di Cauchy di numeri reali; allora $(x_n)_{n \in \mathbb{N}}$ è una successione convergente.*

DIMOSTRAZIONE. Se $(x_n)_{n \in \mathbb{N}}$ è una successione di Cauchy, è limitata. Ha quindi una sottosuccessione estratta convergente ad un limite finito, e quindi tutta la successione, essendo di Cauchy, converge allo stesso limite. $\square$

¹ Un punto $\xi \in \mathbb{R}$ è un *punto di accumulazione* di (o per) $X \subseteq \mathbb{R}$ se ogni intervallo aperto che contiene ξ contiene infiniti punti di X .

² Un sottoinsieme $X \subseteq \mathbb{R}$ è uno *spazio metrico completo* se ogni successione di Cauchy di punti di X è convergente ad un limite che appartiene ad X . Ad esempio, $\mathbb{Q}$ e $(0, 1]$ non sono completi, mentre $\mathbb{R}$ lo è.

Nella costruzione di Dedekind, o nella impostazione assiomatica con la completezza assunta nel senso tecnico della teoria degli insiemi totalmente ordinati (ogni sottoinsieme non-vuoto e superiormente limitato ha estremo superiore), ci si ferma qui. Costruendo invece i reali col modello di Méray–Cantor, per iniziare l’Analisi matematica dobbiamo provare che.

La completezza metrica di $\mathbb{R}$ implica l’esistenza dell’estremo superiore:

Teorema 5.2.6 (Esistenza dell’estremo superiore) *Sia X un sottoinsieme di $\mathbb{R}$ non-vuoto e superiormente limitato; allora esiste $\xi = \sup(X)$.*

DIMOSTRAZIONE. Si procede esattamente come nel teorema (3.2.15) nella costruzione del modello di Méray–Cantor. $\square$

La Figura 5.1 indica i diversi possibili percorsi didattici iniziali dell’analisi matematica a seconda dei differenti “punti di ingresso”, cioè a seconda della definizione di $\mathbb{R}$ adottata.

5.3 Inizia l’analisi matematica

Un teorema essenziale per costruire l’analisi in $\mathbb{R}$ è il Teorema di Heine–Borel, detto talvolta Teorema di Heine–Pincherle–Borel, aggiungendo ai nomi di Heine e di Borel quello del matematico Salvatore Pincherle (1853–1936). Il Teorema di Heine–Borel verrà utilizzato nel Capitolo 8 per dimostrare il Teorema 8.3.3.

Teorema 5.3.1 (Heine–Borel) *Se $((c_i, d_i))_{i \in \mathcal{I}}$ è un insieme arbitrario di intervalli (c_i, d_i) aperti di $\mathbb{R}$ la cui unione contiene un intervallo chiuso e limitato $[a, b]$, allora esiste un sottoinsieme finito di tali intervalli $(c_{i_1}, d_{i_1}), (c_{i_2}, d_{i_2}), \dots, (c_{i_N}, d_{i_N})$ la cui unione $\cup_{j=1}^N (c_{i_j}, d_{i_j})$ contiene $[a, b]$.*

DIMOSTRAZIONE. Supponiamo per assurdo che non si possa coprire³ $I_0 = [a, b]$ con una sottoinsieme finito di intervalli aperti (c_i, d_i) . Considerando i due intervalli $[a, (a+b)/2]$ e $[(a+b)/2, b]$, almeno uno dei

³ Si dice che una famiglia $(A_i)_{i \in \mathcal{I}}$ di insiemi *copre* un insieme E se $E \subseteq \cup_{i \in \mathcal{I}} A_i$.

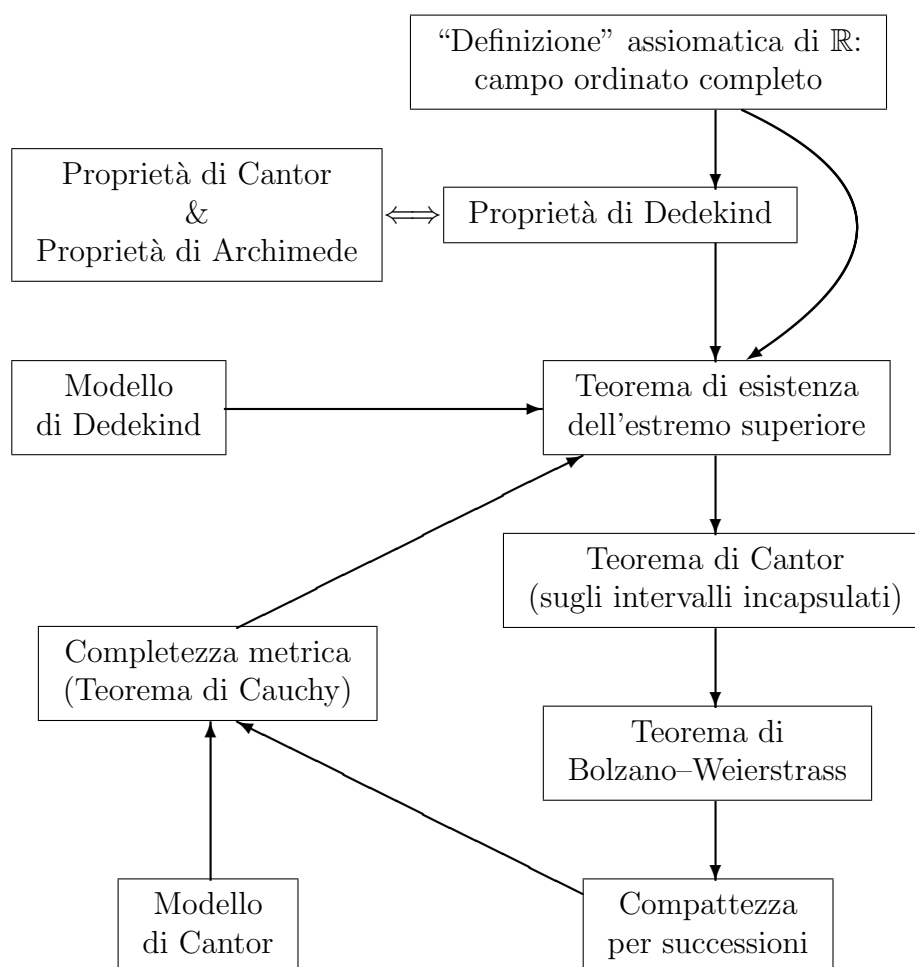


Figura 5.1: Schema dei diversi possibili percorsi didattici iniziali dell'analisi matematica a seconda della definizione di $\mathbb{R}$ adottata.

due non può essere coperto da una sottoinsieme finito di intervalli aperti (c_i, d_i) . Chiamiamo tale intervallo I_1 . Procediamo con tale procedura di bisezione ottenendo una successione $(I_n)_{n \in \mathbb{N}}$ di intervalli chiusi decrescenti per inclusione. Per il Teorema di Cantor 5.2.2 esiste un punto ξ che appartiene a tutti gli I_n . Ora, poiché la unione degli (c_i, d_i) contiene $I_0 = [a, b]$, deve esistere un intervallo $L = (c_k, d_k)$ al quale $\xi \in I_0$ appartiene. Ma, avendo che la lunghezza degli intervalli I_n tende a zero, esiste un intervallo I_m contenuto in L , il che porta ad una contraddizione, in quanto, per costruzione, I_m non può essere coperto da sottoinsieme finito di intervalli aperti (c_i, d_i) (in questo caso $\{L\}$). $\square$

Dal teorema di Heine–Borel segue immediatamente il teorema di Weierstrass sull'esistenza di massimo e minimo per una funzione reale continua definita su un intervallo chiuso e limitato. Dal teorema di Weierstrass segue il teorema di Rolle, e da questo la formula di Taylor, e pian piano tutto il calcolo differenziale.

Capitolo 6

Gli allineamenti di cifre

*Ai modesti o vanitosi ai violenti o timorosi
do, cantando gaio ritmo, logaritmo...*

GIORGIO RABBENO,
 $e = 2.718281828459 \dots$

6.1 Introduzione

In questo capitolo discutiamo l'introduzione di $\mathbb{R}$ con gli allineamenti di cifre decimali. Per semplicità usiamo gli *allineamenti binari* in luogo di quelli decimali. Mostriamo che allineamento di cifre può venir interpretato non tanto come lista dei coefficienti di una serie di potenze, ma piuttosto come *codice* per la descrizione di una sezione di Dedekind, e quindi di un numero reale. Pertanto, la conclusione principale del capitolo è che la costruzione di $\mathbb{R}$ con gli allineamenti decimali è nient'altro che un caso particolare della costruzione di Dedekind, della quale conseguentemente conserva i difetti: le ambiguità del tipo $0.99999\dots = 1.00000\dots$ nella rappresentazione decimale dei reali hanno la loro ineliminabile radice nell'ambiguità della rappresentazione di una sezione di $\mathbb{Q}$ di prima specie, nella quale il numero r “dove si taglia” può indifferentemente essere messo nella classe inferiore o nella classe superiore. Questa “spiegazione” dell'uguaglianza $0.99999\dots = 1.00000\dots$ ci appare molto più convincente della mera constatazione che $\sum_{k=1}^{\infty} 9/10^k = 1$.

Come esercizi, proponiamo le implementazioni in linguaggio R di alcuni sviluppi binari (ad esempio della $\sqrt{2}$ a partire dalla sua definizione come sezione).

6.2 Gli allineamenti binari

Un *alfabeto* $\mathcal{A}$ di $m \geq 2$ *simboli* è una m -pla $(a_1, a_2, \dots, a_m)$ di elementi a_i ($i = 1, \dots, m$) distinti (i simboli), che si assume totalmente ordinata con l'ordinamento indotto da quello naturale degli indici:

$$a_1 < a_2 < a_3 < \dots < a_m$$

Un alfabeto con due simboli è quello *binario*, tipicamente $\mathcal{A} = (0, 1)$. L'alfabeto *italiano* è

$$\mathcal{A} = (a, b, c, d, e, f, g, h, i, l, m, n, o, p, q, r, s, t, u, v, z)$$

ed ha $m = 21$ simboli; l'ordinamento è quello “alfabetico”. L'alfabeto *genetico* è $\mathcal{A} = (A, T, C, G)$, dove i simboli rappresentano le basi del DNA, e precisamente A = adenina, T = timina, G = guanina, C = citosina. L'alfabeto *decimale* è quello delle $m = 10$ cifre decimali $\mathcal{A} = (0, 1, 2, 3, 4, 5, 6, 7, 8, 9)$.

Una *parola* s di lunghezza n tratta dall'alfabeto $\mathcal{A}$ è una n -pla $s \in \mathcal{A}^n$, ossia una n -pla (ordinata)

$$s = (s_1, s_2, \dots, s_n), \quad \text{dove } s_k \in \mathcal{A}, (k = 1, 2, \dots, n)$$

di simboli di $\mathcal{A}$. Una parola $s = (s_1, s_2, \dots, s_n)$ si scrive tipicamente scrivendo i simboli uno dopo l'altro, cioè nella forma compattata $s = “s_1 s_2 \dots s_n”$.

Diremo invece *allineamento* tratto dall'alfabeto $\mathcal{A}$ una *successione*

$$s = (s_1, s_2, s_3, \dots), \quad \text{dove } s_k \in \mathcal{A}, (k = 1, 2, \dots)$$

di simboli di $\mathcal{A}$, ovvero un elemento $s \in \mathcal{A}^{\mathbb{N}^*}$, dove $\mathbb{N}^* = \mathbb{N} \setminus \{0\}$. Tipicamente gli allineamenti si rappresentano come le parole, ossia “allineando” i vari simboli uno dopo l'altro nell'ordine assegnato:¹

$$s = s_1 s_2 s_3 \dots$$

¹ Naturalmente non possiamo scriverli tutti: i tre puntini (...) rappresentano quelli non scritti.

Per esempio

$$s = 891834573432835947800061913789 \dots$$

rappresenta un *allineamento decimale*, mentre

$$s = 001000111000100101101010110000 \dots$$

rappresenta un *allineamento binario*. Ci interessano qui in particolare gli allineamenti binari.

6.2.1 La rappresentazione binaria

Un allineamento binario rappresenta “filosoficamente” una successione di infinite (numerabili) scelte fra due possibilità; ad esempio, i risultati di una successione (infinita) di lanci di una moneta con esiti *testa* = 0 oppure *croce* = 1 possono essere descritti da un allineamento binario. Mostriamo ora come un allineamento binario possa offrire una descrizione “economica” di un numero reale.

Abbiamo già visto, nel corso della definizione delle operazioni algebriche nel modello di Dedekind, che la rappresentazione di un numero reale $x = (A, B)$ con entrambe le classi è sovrabbondante: infatti è sufficiente assegnare la classe inferiore, nella forma di una semiretta sinistra di Russell aperta a destra, per identificare la sezione. Vedremo ora che in un certo senso anche le semirette di Russell (le classi inferiori) sono sovrabbondanti: basterà assegnare un’opportuna successione di elementi della classe inferiore per identificarla perfettamente.

Consideriamo quindi una sezione di Dedekind $x = (A, B)$ del campo razionale $\mathbb{Q}$, ossia un numero reale; supponiamo per semplicità che in $\mathbb{R}$ sia $x > 0$. La classe inferiore A contiene quindi sia lo 0 sia razionali positivi, cioè $A^+ = A \cap \mathbb{Q}^+ \neq \emptyset$. La classe A è superiormente limitata da un qualunque elemento di B e pertanto l’intersezione $A \cap \mathbb{N}$ è finita ed ha un massimo elemento $a_0 = d$, che diremo il *basamento* di x . Quindi $b_0 = d + 1 \in B$. Cerchiamo di identificare gli elementi di A nell’intervallo $[a_0, b_0) = [d, d + 1)$. Abbiamo già incontrato in precedenza (Teoremi 2.7.4, 3.2.15 e 5.2.3) il seguente algoritmo di bisezione.

- Consideriamo il punto di mezzo $t_0 = (a_0 + b_0)/2$: definiamo un intervallo

$$[a_1, b_1) = \begin{cases} [a_0, t_0) & \text{se } t_0 \in B, \text{ e poniamo } s_1 = 0 \\ [t_0, b_0) & \text{se } t_0 \in A, \text{ e poniamo } s_1 = 1 \end{cases}$$

Si osservi che $a_1 = d + s_1 \cdot 2^{-1}$.

- Consideriamo il punto di mezzo $t_1 = (a_1 + b_1)/2$: definiamo un intervallo

$$[a_2, b_2) = \begin{cases} [a_1, t_1) & \text{se } t_1 \in B, \text{ e poniamo } s_2 = 0 \\ [t_1, b_1) & \text{se } t_1 \in A, \text{ e poniamo } s_2 = 1 \end{cases}$$

Si osservi che $a_2 = a_1 + s_2 \cdot 2^{-2} = d + s_1 \cdot 2^{-1} + s_2 \cdot 2^{-2}$.

- ...

Si osservi ancora che l'intervallo generico $[a_n, b_n)$ ha ampiezza 2^{-n} , e gli estremi sono

$$a_n = d + \sum_{k=1}^n s_k \cdot 2^{-k} \in A$$

e $b_n \in B$. Ora la successione $(a_n)_{n \geq 0}$ identifica la classe inferiore A , nel senso che saturando inferiormente i termini della successione $(a_n)_{n \geq 0}$ si ottiene A . Basta provare che un qualsiasi $a \in A$ è superato da qualche a_n della successione. Poiché A non ha massimo, esiste $a' \in A$, $a' > a$; calcoliamo la differenza $a' - a > 0$ e sia n tale che $2^{-n} < a' - a$ (proprietà di Archimede² in $\mathbb{Q}$). Se fosse $a_n \leq a$, avremmo $b_n = a_n + 2^{-n} < a + (a' - a) = a'$, assurdo in quanto $b_n \in B$.

D'altra parte, la successione $(a_n)_{n \geq 0}$ è a sua volta perfettamente identificata dall'intero $d \geq 0$ e dall'allineamento binario $s_1 s_2 s_3 \dots$; quindi la coppia $(d, s_1 s_2 s_3 \dots)$ identifica perfettamente A , quindi la sezione (A, B) , quindi il numero reale x . Si scrive comunemente

$$x := d . s_1 s_2 s_3 \dots$$

Abbiamo così collegato la definizione di numero reale alla Dedekind con la *rappresentazione binaria* del numero stesso. È importante osservare

² Siccome $2^n = (1+1)^n = 1+n+\binom{n}{2}+\dots > 1+n$, basta prendere $n+1 > 1/(a'-a)$ per avere $1/(a'-a) < 2^n$, e quindi $a'-a < 2^{-n}$.

che la scrittura $x := d.s_1s_2s_3\dots$ è una scrittura riassuntiva per indicare la definizione di $x := A$ attraverso la sua classe inferiore.

Naturalmente, se in luogo dell'alfabeto binario si intende utilizzare l'alfabeto decimale, basta suddividere l'intervallo $[a_0, b_0) = [d, d+1)$ ed i sottointervalli successivi in 10 parti invece che in 2, e assegnare le cifre decimali di conseguenza. Per comprendere la natura della rappresentazione decimale è però meglio riferirsi alla rappresentazione binaria.³

Conviene ancora osservare che, nel metodo che abbiamo descritto, è *impossibile* che nella rappresentazione binaria di un numero reale $x := d.s_1s_2s_3\dots$ si trovino da un certo indice j in poi tutte le cifre binarie uguali a 0. Questo infatti implicherebbe che da quell'indice in poi (ovvero per $n \geq j$) la successione $(a_n)_{n \geq 0}$ sarebbe costante (ovvero sarebbe $a_n = a_j$). Allora tutti i razionali $r > a_j$ sarebbero in B , e quindi A avrebbe massimo $= a_j$. Possiamo ora reinterpretare i tre casi che si presentano per una sezione originale di Dedekind $x = (A, B)$, ossia quella definita soltanto dalle (S0) ed (S1) della Definizione 4.2.1:

- (i) A non ha massimo e B ha minimo r ;
- (ii) A non ha massimo e B non ha minimo;
- (iii) A ha massimo r e B non ha minimo;

Se si aggiunge, come noi abbiamo fatto, la condizione

(S2) A non ha massimo

il caso (iii) è escluso, il caso (ii) corrisponde agli irrazionali, il caso (i) ai razionali con rappresentazione binaria *non 0-periodica*; se fosse stata aggiunta in luogo di (S2) la condizione simmetrica

(S2bis) B non ha minimo

il caso (i) sarebbe escluso, il caso (ii) sarebbe stato ugualmente corrispondente agli irrazionali, ed il caso (iii) ai razionali con rappresentazione binaria *non 1-periodica*. Se si assumono solo (S0) ed (S1), allora ripetendo

³ Quando sia necessario distinguere nello stesso contesto le rappresentazioni di un numero reale in basi diverse, $b = 2$ oppure $b = 10$, apporremo un pedice 2 alla rappresentazione binaria. Ad esempio, $9 = 101_2$.

il ragionamento per bisezione sono possibili sia la rappresentazione binaria 0-periodica che quella 1-periodica. Ancora una volta nel caso binario le cose sono semplici più che nel caso decimale, dove lo stesso numero razionale può ammettere due rappresentazioni, sia una 0-periodica che una 9-periodica. La cosa didatticamente (e matematicamente) errata è di ritenere⁴ la rappresentazione decimale 0-periodica una rappresentazione *limitata*, mentre la rappresentazione decimale 9-periodica viene ritenuta una rappresentazione *illimitata*: questa asimmetria non esiste. La rappresentazione come successione di scelte non può essere limitata, perché le scelte devono essere “tutte” effettuate. Per chiarire ancora il punto con un esempio, prendiamo $x = 2$, o meglio $x = \hat{2}$, la cui classe inferiore, nella teoria qui presentata con la aggiunta di (S2), è $A = \{a \in \mathbb{Q} \mid a < 2\}$ (classe priva di massimo). Quindi $A \cap \mathbb{N} = \{0, 1\}$ ha un massimo elemento $a_0 = d = 1$, che sarebbe il *basamento* di $\hat{2}$. Si osservi che $2 = b_0 \in B$. Abbiamo $t_0 = (a_0 + b_0)/2 = 3/2 \in A$, quindi $s_1 = 1$; abbiamo $a_1 = t_0 = 1 + 1/2$, $b_1 = b_0 = 2$, $t_1 = (a_1 + b_1)/2 = 1 + 1/2 + 1/4 \in A$, quindi $s_2 = 1$; ecc. In definitiva la nostra rappresentazione per $x = \hat{2}$ è un po' strana:

$$\hat{2} = 1.1111111111111111 \dots$$

ma è quella che riassume la classe inferiore di 2 priva di massimo; ossia, saturando inferiormente i numeri $1, 1 + \frac{1}{2}, 1 + \frac{1}{2} + \frac{1}{4}, \dots$, si trovano *tutti* i numeri razionali < 2 .

La cosiddetta “ambiguità” nella rappresentazione binaria o decimale che sia, ha quindi la sua ineliminabile radice nell’ambiguità della rappresentazione di una sezione di $\mathbb{Q}$ di prima specie, con il numero r “dove si taglia” che può indifferentemente essere messo nella classe superiore (come abbiamo fatto noi) o nella classe inferiore.

È ancora importante notare che la rappresentazione binaria dei reali che qui abbiamo presentato si sviluppa esclusivamente in $\mathbb{Q}$.

Topologia degli allineamenti binari

Consideriamo l’insieme $E = \Sigma_2$ di tutti gli allineamenti binari a di due simboli, zero ‘0’ e uno ‘1’, con indice $k \geq 1$:

$$a = (a_1, a_2, a_3, a_4, a_5, a_6, \dots) = “a_1 a_2 a_3 a_4 a_5 a_6 \dots”$$

⁴ Come in una certa tradizione scolastica.

con esclusione di quelle definitivamente uguali a '1'. Con la distanza $d(a, b) = \sum_{k=1}^{\infty} |a_k - b_k| 2^{-k}$, E diventa uno spazio metrico limitato. È immediato riconoscere che ogni allineamento $a \in \Sigma_2$ può essere identificata con un numero reale $\nu(a) = \sum_{k=1}^{\infty} a_k 2^{-k} \in [0, 1)$. L'applicazione $\nu : \Sigma_2 \rightarrow \mathbb{R}$ è continua, anzi, è lipschitziana di costante 1:

$$|\nu(a) - \nu(b)| = \left| \sum_{k=1}^{\infty} (a_k - b_k) 2^{-k} \right| \leq \sum_{k=1}^{\infty} |a_k - b_k| 2^{-k} = d(a, b)$$

Per esempio, se scegliamo $a = "11100000 \dots"$ ($a_1 = a_2 = a_3 = 1$ e $a_k = 0$ per $k \geq 4$) e $b = "10110000 \dots"$ ($b_1 = 1, b_2 = 0, b_3 = b_4 = 1$ e $b_k = 0$ per $k \geq 5$), abbiamo

$$\nu(a) = \frac{1}{2} + \frac{1}{4} + \frac{1}{8}, \quad \nu(b) = \frac{1}{2} + \frac{1}{8} + \frac{1}{16}, \quad |\nu(a) - \nu(b)| = \frac{1}{4} - \frac{1}{16} = \frac{3}{16}$$

mentre

$$d(a, b) = \sum_{k=1}^{\infty} |a_k - b_k| 2^{-k} = \frac{0}{2} + \frac{1}{4} + \frac{0}{8} + \frac{1}{8} = \frac{3}{8}$$

Avendo escluso gli allineamenti definitivamente uguali a '1', la $\nu : \Sigma_2 \rightarrow [0, 1)$ è biiettiva, e la inversa ν^{-1} è la rappresentazione binaria non 1-periodica di un numero reale dell'intervallo $[0, 1)$. È importante notare che l'origine dei problemi tecnici nell'introduzione dei reali come allineamenti di cifre è il fatto che ν non è un omeomorfismo, in quanto la rappresentazione binaria $\nu^{-1} : [0, 1) \rightarrow \Sigma_2$ non è continua: infatti, posto

$$\theta = 1/2 = .1_2, \quad \theta_n = 1/2 - 1/2^n = \underbrace{.011111 \dots 11}_n_2$$

si ha che $\theta_n \rightarrow \theta$ in $[0, 1)$, mentre in Σ_2

$$d(\underbrace{"100000 \dots 00"}_n, \underbrace{"011111 \dots 11"}_n) = \sum_{k=1}^n 1 \cdot 2^{-k} \rightarrow 1$$

6.2.2 Implementazione in R

Vediamo alcuni esempi di rappresentazione binaria implementata in R. Vediamo uno pseudo-codice per l'algoritmo di rappresentazione binaria del numero reale $x = (A, B)$.

```

1.  INPUT:  $N \geq 1$ 
2.  calcolo:  $d = \max(\mathbb{N} \cap A)$ 
3.  inizializzo:  $a = d, b = d + 1$ 
4.  inizializzo la lista delle cifre:  $s = \emptyset$ 
5.      FOR  $n = 1, 2, \dots, N$ 
6.          calcolo il punto di mezzo:  $t = (a + b)/2$ 
7.              IF  $t \in A$ 
8.                  THEN  $a = t$ , aggiungo la cifra 1 ad  $s$ 
9.                  ELSE  $b = t$ , aggiungo la cifra 0 ad  $s$ 
10.             END IF
11.     END FOR
12.  OUTPUT:  $s$ 

```

Per implementare il codice in R useremo una funzione logica `test` di una variabile razionale q per stabilire se $q \in A$ (`test(q) = TRUE`) o se $q \notin A$ (`test(q) = FALSE`). Negli esempi successivi calcoleremo le prime $N = 30$ cifre binarie del numero dato, e, per semplicità, introdurremo direttamente il valore del basamento. Il codice in R risulta quindi:

```

1. N <- 30
2. d <- 1
3. a <- d; b <- d+1
4. s <- c()
5.   for (j in (1:N)){
6.     t <- (a+b)/2;
7.     if (test(t))
8.       {a <- t; s <- c(s,1)}
9.     else {b <- t; s <- c(s,0)}
10.   }
11.   #
12. print(paste(x =,d,,paste(s,collapse=''),...),quote=F)

```

Rappresentazione binaria di $x = \widehat{2}$

La classe inferiore di $x = \widehat{2}$ è $A = \{q \in \mathbb{Q} \mid q < 2\}$. La funzione logica `test` è

```
test <- function(q){q < 2}
```

```
print(paste("x =",d,".",paste(s,collapse=""),"..."),quote=F)
[1] x = 1 . 111111111111111111111111111111111111 ...
```

$$x = d + \sum_{i=1}^N s_j 2^{-j} = d + s_1 \frac{1}{2} + s_2 \frac{1}{4} + s_3 \frac{1}{8} + s_4 \frac{1}{16} + \dots + s_{30} \frac{1}{2^N}$$

```
x <- d + as.real(s %*% 2^(-(1:N))); x
[1] 2
```

Rappresentazione binaria di $x = \sqrt{2}$

```
test <- function(q){q^2 < 2}
```

[1] $x = 1 . 011010100000100111100110011001 \dots$

```
[1] 1.414214
```

```
test <- function(q){q < 3/5}
```

[1] $x = 0.100110011001100110011001100110 \dots$

[1] 0.6

Rappresentazione binaria di $\phi = \frac{1 + \sqrt{5}}{2}$ (rapporto aureo)

La funzione logica `test` è

```
test <- function(q){(2*q - 1)^2 < 5}
```

Il basamento è $d = 1$. L'output è:

```
[1] x = 1 . 100111100011011101111001101110 ...
```

Il valore decimale è

```
[1] 1.618034
```

Osservazione 6.2.1 La rappresentazione binaria o decimale come presentata funziona anche per numeri $x = (A, B)$ negativi. Se $x < 0$ il basamento d di x è il massimo intero in A , ossia $d = \max(\mathbb{Z} \cap A)$. Ad esempio per $x = \widehat{-11/4}$ avremmo un basamento $d = -3$; siccome A si “spinge a destra” di $1/4$; la rappresentazione decimale sarebbe allora $-11/4 = (-3).25$ cioè $1/4$ a destra di -3 .

In pratica, come sappiamo, non si usa questa rappresentazione, ma si preferisce aggiungere 1 bit alla rappresentazione dei numeri positivi, calcolando la rappresentazione binaria o decimale di $|x|$ ed aggiungendovi all'inizio un simbolo $+$ se $x > 0$ oppure un simbolo $-$ se $x < 0$. Così la rappresentazione “accettata” di $-11/4$ non è $(-3).25$ ma -2.75 .

Esempio: rappresentazione binaria di $x = -11/4$

La funzione logica `test` è

```
test <- function(q){q < -11/4}
```

Il basamento d è il massimo intero $< -11/4$, ossia $d = -3$. L'output è:

```
[1] x = -3 . 001111111111111111111111111111 ...
```

Il valore decimale è

```
[1] -2.75
```

6.2.3 Implementazione in *Mathematica*

Se utilizziamo un Computer Algebra System come *Mathematica*, possiamo ottenere risultati maggiormente informativi. Vediamo come esempio la rappresentazione binaria $x = \sqrt{2}$. Se vogliamo che questa corrisponda ad una rappresentazione decimale con, ad esempio, 25 cifre decimali “dopo la virgola”, ci serviranno $N = \log_2 10^{25} \approx 83.0482$ cifre binarie: facciamo quindi $N = 84$. Definiamo in *Mathematica* la funzione `test`:

```
In[1]:= test[q_]:= q^2<2
In[2]:= {test[1], test[2]}
Out[2]= {True, False}
```

Procediamo come in precedenza:

```
In[3]:=
n=84; d=1; a=d; b=d+1; s=""
Do[ t=(a+b)/2;
  If [test[t],
    {a=t; s=s<>"1"},
    {b=t; s=s<>"0"}],
  {j,1,n}]
Print[ToString[d]<>" . "<>s]
x=d+ToExpression[Characters[s]] . Table[2^-k,{k,1,n}]

Out[3]= 1 . 01101010000010011110011001100111111100111011
        1100110010010000100010110010111110110001

Out[6]= 27354868640032294882193329
        19342813113834066795298816
```

Si noti che *Mathematica* ha fornito una rappresentazione razionale di

$$a_{84} = \frac{27354868640032294882193329}{19342813113834066795298816}$$

Abbiamo quindi trovato nella classe inferiore A della sezione $(A, B) = \sqrt{2}$ un numero razionale, ossia a_{84} , che dista “pochissimo” da elementi della classe superiore B , precisamente meno di

$$2^{-84} = \frac{1}{19342813113834066795298816} \approx 5.1698788 \times 10^{-26}$$

Poiché il denominatore di a_{84} è una potenza di 2, la rappresentazione decimale di a_{84} termina necessariamente con $\dots 25$ e 0 periodico.⁵ Si ha precisamente (ed *esattamente*):

$$a_{84} = 1.414213562373095048801688713217188396575733390359008723180522792972624301910400390625\overline{0}$$

Per essere pedanti, proviamo che $\sqrt{2} - \widehat{a_{84}} \leq \widehat{2^{-84}}$, ossia (siamo in un corpo ordinato!) $\sqrt{2} \leq \widehat{a_{84}} + \widehat{2^{-84}} = \widehat{b_{84}}$. La classe inferiore di $\sqrt{2}$ è A ; la classe inferiore di $\widehat{b_{84}}$ è $C = \{q \in \mathbb{Q} \mid q < b_{84}\}$, ed evidentissimamente $A \subseteq C$. Per questi motivi si usa dire che

$$\sqrt{2} \approx 1.4142135623730950488016887$$

con 25 cifre decimali esatte dopo la virgola.

6.3 Costruzione di $\mathbb{R}$ tramite allineamenti

Si comprende quindi come sia possibile costruire $\mathbb{R}$ tramite allineamenti binari o decimali. Una coppia del tipo $x = (n, s)$ dove $n \in \mathbb{N}$ ed s è un allineamento binario non 0-periodico identifica la classe inferiore di una sezione (A, B) di $\mathbb{Q}$ (che verifica (S0)–(S1)–(S2)). Basta quindi ripetere la costruzione di Dedekind “lavorando” con tali allineamenti invece che con le sezioni o con le loro classi inferiori. I problemi tecnici sorgono quando si vogliono introdurre le operazioni di somma e prodotto, e le difficoltà sono legate essenzialmente al fatto già in precedenza discusso che la rappresentazione binaria ν^{-1} non è continua. In pratica, se si sommano i valori approssimati, fino ad una certa cifra, di due reali, la loro somma o il loro prodotto non è l'approssimazione, fino alla stessa cifra, della somma o del prodotto dei due numeri. Per esempio si ha: $7/9 + 5/9 = 4/3$, ovvero, in termini di allineamenti decimali $0.77777\dots + 0.55555\dots = 1.33333\dots$. Se si usano le approssimazioni fino alla quarta cifra, per esempio per troncamento, si ottiene invece: $0.7777 + 0.5555 = 1.3332$ che è corretto solo fino alla terza cifra. E il discorso non cambia se si usano le approssimazioni per arrotondamento: si ottiene $0.7778 + 0.5556 = 1.3334$.

⁵ È quindi quello che la tradizione scolastica cui abbiamo precedentemente accennato chiama un numero decimale “limitato”.

Le difficoltà si superano appunto seguendo la costruzione di Dedekind, dimostrando prima l'esistenza dell'estremo superiore, e poi definendo le operazioni tramite questo. Ad esempio si definisce la somma $\sqrt{2} + \sqrt{3}$ come l'estremo superiore delle somme di troncamenti della rappresentazione decimale di $\sqrt{2}$ con troncamenti della rappresentazione decimale di $\sqrt{3}$. Analogamente si definisce il prodotto $\sqrt{2} \cdot \sqrt{3}$ come l'estremo superiore dei prodotti di troncamenti della rappresentazione decimale di $\sqrt{2}$ con troncamenti della rappresentazione decimale di $\sqrt{3}$.⁶ Per una realizzazione concreta di questo processo didattico, vedi G. C. Barozzi [4].

6.4 Costruzione di $\mathbb{R}$ tramite classi contigue

Un'altra costruzione di $\mathbb{R}$ è quella tramite *classi contigue*. Abbiamo visto nella Sezione 6.2.1 sugli allineamenti binari che un numero reale alla Dedekind $x = (A, B)$ viene identificato da un'opportuna successione $(a_n)_{n \geq 0}$, o, il che è lo stesso, da un'opportuna coppia di successioni $((a_n)_{n \geq 0}, (b_n)_{n \geq 0})$ per le quali è $b_n = a_n + 2^{-n}$ per ogni $n \geq 0$. Conseguentemente $(\{a_n\}_{n \geq 0}, \{b_n\}_{n \geq 0})$ è una particolare coppia di classi contigue di numeri razionali. Si comprende quindi come sia possibile costruire $\mathbb{R}$ tramite generiche coppie di classi contigue di numeri razionali, anche qui ripetendo la costruzione di Dedekind “lavorando” con tali coppie di classi contigue di numeri razionali invece che con le sezioni. La realizzazione concreta di questo processo didattico ha avuto un certo successo in Italia, forse anche in seguito al suo utilizzo nelle *Lezioni di algebra complementare* di A. Capelli [13].

⁶ Questo esempio è critico, in quanto, per “scoprire” alla fine che $\sqrt{2}\sqrt{3} = \sqrt{6}$ occorre dimostrare che per la moltiplicazione $\cdot$ così definita valgono le proprietà associative e commutativa: in tal caso infatti si ha $(\sqrt{2}\sqrt{3})^2 = (\sqrt{2}\sqrt{3})(\sqrt{2}\sqrt{3}) = (\sqrt{2}\sqrt{2})(\sqrt{2}\sqrt{3}) = 6$, dal che segue che $\sqrt{2}\sqrt{3} = \sqrt{6}$.

Capitolo 7

Gli iperreali

M. Jourdain.— *Quoi, quand je dis: “Nicole, apportez-moi mes pantoufles, et me donnez mon bonnet de nuit”, c’est de la prose?*

Maître de philosophie.— *Oui, Monsieur.*

M. Jourdain.— *Par ma foi, il y a plus de quarante ans que je dis de la prose, sans que j’en susse rien; et je vous suis le plus obligé du monde, de m’avoir appris cela.^a*

MOLIÈRE,

Le Bourgeois gentilhomme

^a M. Jourdain.— *Come? quando dico: “Nicoletta, portami le pantofole, e dammi il berretto da notte”, è prosa?* Maestro di filosofia.— *Sì, signore.*
M. Jourdain.— *Per tutti i diavoli! Sono più di quarant’anni che parlo in prosa senza saperlo. Vi sono molto grato di avermi informato.*

7.1 Introduzione

Il successo straordinario delle costruzioni classiche dei numeri reali consiste essenzialmente nella possibilità di definire la nozione di limite in un modo che oggi riteniamo “corretto”. Ad esempio, quando scriviamo la definizione di derivata

$$\lim_{\epsilon \rightarrow 0} \frac{f(a + \epsilon) - f(a)}{\epsilon} = f'(a) \quad (7.1)$$

utilizziamo un “operatore”, ossia $\lim_{\epsilon \rightarrow 0}$, che ci fa passare da un’espressione che contiene esplicitamente ϵ ad una che non lo contiene più. In un certo senso quindi, le costruzioni classiche di $\mathbb{R}$ forniscono una risposta al problema posto dal filosofo George Berkeley (1685–1753), che, nel suo trattato *The Analyst: or a discourse addressed to an infidel mathematician*, aveva attaccato duramente il calcolo infinitesimale di Newton e Leibniz osservando che qualche volta gli “infinitesimi” venivano e dovevano essere considerati diversi da zero (come quando ϵ compare al denominatore del rapporto incrementale), mentre altre volte diventavano uguali a zero (come quando ϵ scompare nell’espressione di $f'(a)$). È forse in questa risposta il senso profondo della *arimetizzazione dell’analisi*, ossia la fondazione formale di $\mathbb{R}$ sui razionali (cf. [23], p. 10).

Tuttavia, è anche implicitamente o esplicitamente ammesso da chi insegna o studia matematica che gli ϵ ed i δ della definizione classica di limite costituiscono un importante problema didattico. Per questo motivo inseriamo in questo testo un Capitolo dedicato alla risposta alternativa data alle obiezioni di Berkeley da Abraham Robinson, un logico matematico ideatore della cosiddetta *analisi non-standard* [51].

L’analisi non-standard costruisce infatti un corpo ordinato $\mathbb{K}$, necessariamente non archimedeo, i cui elementi sono detti *numeri iperreali*, e che contiene sia i numeri reali “standard” di $\mathbb{R}$ che numeri *infinitesimi*, ossia elementi $\epsilon \in \mathbb{K}$ tali che per ogni reale “standard” $x \in \mathbb{R}$, $x > 0$, sia $|\epsilon| < x$. Quello che ha consentito un certo successo didattico dell’analisi non-standard (soprattutto negli Stati Uniti; per l’Italia si vedano ad esempio [7] e [43]) è il fatto che, per gli scopi dell’insegnamento e della ricerca in analisi matematica, alla fine non è necessario studiare preliminarmente decine di pagine di logica matematica, come in Robinson [51].

Abbiamo infatti aperto il capitolo con il celebre passo del *Bourgeois gentilhomme* di Molière, nel quale il protagonista, Monsieur Jourdain, scopre di parlare in prosa da più di quarant’anni senza saperlo, per enfatizzare che, come vedremo, la costruzione del corpo dei *numeri iperreali* rientra perfettamente, e abbastanza facilmente, nello schema generale dell’arimetizzazione: la costruzione dei numeri iperreali ha infatti la medesima struttura della costruzione dei numeri reali di Méray–Cantor.

7.2 Costruzione degli iperreali

La costruzione dei numeri iperreali avviene a partire dai reali, e si sviluppa secondo le stesse idee della costruzione dei numeri reali a partire dai razionali con il metodo di Méray–Cantor. Ricordiamo che il metodo di Méray–Cantor parte dall’anello delle successioni di Cauchy di razionali, dichiara “zero” le successioni razionali convergenti a 0 ed arriva ad $\mathbb{R}$.

Per costruire gli iperreali, in luogo di considerare le sole successioni di Cauchy di numeri razionali, consideriamo l’insieme $\mathbb{R}^{\mathbb{N}}$ di tutte le successioni $p = (p_0, p_1, p_2, \dots) = (p_k)_{k \geq 0}$ di numeri reali. In $\mathbb{R}^{\mathbb{N}}$ definiamo l’addizione e la moltiplicazione per componenti: dati $p = (p_0, p_1, p_2, \dots)$ e $q = (q_0, q_1, q_2, \dots)$, poniamo

$$p + q = (p_0 + q_0, p_1 + q_1, p_2 + q_2, \dots)$$

$$p \cdot q = (p_0 q_0, p_1 q_1, p_2 q_2, \dots)$$

Per ogni $r \in \mathbb{R}$ indichiamo con (r) la successione costante $(r) = (r, r, r, \dots)$. È facilissimo verificare che la struttura algebrica $(\mathbb{R}^{\mathbb{N}}, +, \cdot)$ è un *anello commutativo unitario*, con elemento neutro per la addizione la successione (0) , e con elemento neutro per la moltiplicazione la successione (1) . L’idea sottostante alla costruzione dei numeri iperreali è quella di utilizzare la struttura algebrica $(\mathbb{R}^{\mathbb{N}}, +, \cdot)$, definendo “zero” quelle successioni $(p_k)_{k \geq 0}$ di numeri reali che, in un senso opportuno, sono 0 “all’infinito”. A tale scopo, introduciamo le nozioni di *filtro* e di *ultrafiltro* su $\mathbb{N}$.

7.2.1 Filtri e ultrafiltri

Definizione 7.2.1 *Un filtro su $\mathbb{N}$ è una famiglia non vuota $\mathcal{F}$ di sottoinsiemi di $\mathbb{N}$ che verifica le proprietà seguenti:*

- (F1) *se $A \in \mathcal{F}$, e $A \subseteq B$, allora $B \in \mathcal{F}$;*
- (F2) *se $A \in \mathcal{F}$, $B \in \mathcal{F}$, allora $A \cap B \in \mathcal{F}$;*
- (F3) *$\emptyset$ non appartiene a $\mathcal{F}$.*

Un esempio importante di filtro è dato dal seguente:

Esempio 7.2.2 La famiglia $\mathcal{F}_0$ dei sottoinsiemi F di $\mathbb{N}$ il cui complementare $\mathcal{C}F = \mathbb{N} \setminus F$ è finito, è un filtro su $\mathbb{N}$, detto il filtro di Fréchet.

Le (F1)–(F2)–(F3) si verificano immediatamente. Si osservi che un insieme F di indici appartiene al filtro $\mathcal{F}_0$ di Fréchet se e solo se F contiene una “coda” $\{n \in \mathbb{N} \mid n \geq \bar{n}\}$ di $\mathbb{N}$. L’insieme dei filtri su $\mathbb{N}$ può essere ordinato parzialmente per inclusione. Gli elementi massimali in questo insieme parzialmente ordinato sono detti *ultrafiltri*.

Definizione 7.2.3 Un ultrafiltro su $\mathbb{N}$ è un filtro $\mathcal{F}$ su $\mathbb{N}$ che è massimale per inclusione, ossia tale che, se $\mathcal{G}$ è un filtro su $\mathbb{N}$ che contiene $\mathcal{F}$, allora $\mathcal{G} = \mathcal{F}$.

La proposizione seguente caratterizza gli ultrafiltri.

Proposizione 7.2.4 Un filtro $\mathcal{F}$ su $\mathbb{N}$ è un ultrafiltro se e solo se verifica la proprietà:

(F4) *quale che sia un sottoinsieme $A \subseteq \mathbb{N}$, se $A \notin \mathcal{F}$, allora $\mathcal{C}A \in \mathcal{F}$.*

DIMOSTRAZIONE. Supponiamo che $\mathcal{F}$ sia un filtro che verifica la (F4), e sia $\mathcal{G}$ un filtro tale che $\mathcal{G} \supseteq \mathcal{F}$. Se l’inclusione è stretta, esiste $B \in \mathcal{G}$, $B \notin \mathcal{F}$. Allora $\mathcal{C}B \in \mathcal{F} \subseteq \mathcal{G}$, così $\emptyset = \mathcal{C}B \cap B \in \mathcal{G}$, contraddicendo la (F3) della definizione. Dunque non esiste un filtro $\mathcal{G}$ che contiene strettamente $\mathcal{F}$, cioè $\mathcal{F}$ è un ultrafiltro.

Viceversa, supponiamo che $\mathcal{F}$ sia un ultrafiltro, e sia $A \subseteq \mathbb{N}$ tale che $A \notin \mathcal{F}$. Consideriamo la famiglia

$$\mathcal{G} = \{X \subseteq \mathbb{N} \mid \exists F \in \mathcal{F}, A \cap F \subseteq X\}$$

Risulta $\mathcal{F} \subseteq \mathcal{G}$ (infatti se $F \in \mathcal{F}$ allora $F \subseteq \mathbb{N}$ e $F \supseteq A \cap F$, cioè $F \in \mathcal{G}$) ed è $\mathcal{F} \neq \mathcal{G}$ (ad esempio, $A \in \mathcal{G}$, perché $A \subseteq \mathbb{N}$ e $A \supseteq A \cap F$ per ogni $F \in \mathcal{F}$). Dunque $\mathcal{G}$ non è un filtro, dato che $\mathcal{F}$ è un ultrafiltro, cioè un filtro massimale per inclusione. Osserviamo però che $\mathcal{G}$ non è vuoto e che $\mathcal{G}$ verifica la (F1): infatti, se $B \in \mathcal{G}$ e $D \supseteq B$ si ha immediatamente che $D \in \mathcal{G}$. Inoltre, $\mathcal{G}$ verifica la (F2): infatti, se $B, C \in \mathcal{G}$, allora $B \supseteq A \cap F_0$ e $C \supseteq A \cap F_1$ dove $F_0, F_1 \in \mathcal{F}$ e da questo segue facilmente che $B \cap C \supseteq A \cap (F_0 \cap F_1)$ (si ricordi che $F_0 \cap F_1 \in \mathcal{F}$). Dunque, dato

che $\mathcal{G}$ non è un filtro, deve essere $\emptyset \in \mathcal{G}$. Allora esiste $F \in \mathcal{F}$ tale che $A \cap F = \emptyset$, da cui $F \subseteq \mathcal{C}A$. Da questo segue che $\mathcal{C}A \in \mathcal{F}$ per la (F1) della definizione. $\square$

Vediamo ora un esempio esplicito di ultrafiltro.

Esempio 7.2.5 *La famiglia $\mathcal{U}_m$ dei sottoinsiemi U di $\mathbb{N}$ ai quali appartiene un fissato elemento $m \in \mathbb{N}$ è un ultrafiltro, detto ultrafiltro principale.*

Le proprietà (F1)–(F4) sono verificate banalmente. $\square$

7.2.2 Il Lemma di Zorn

Gli ultrafiltri principali sono in un certo senso “banali”: la loro esistenza segue semplicemente e direttamente dalla loro definizione. Si pone quindi in modo naturale il problema di sapere se esistono o meno ultrafiltri su $\mathbb{N}$ che *non siano principali*. In effetti, dimostreremo nel Teorema 7.2.7 che *esistono* ultrafiltri non principali su $\mathbb{N}$, ma la dimostrazione non consisterà nella esibizione esplicita di un esempio, come nel caso degli ultrafiltri principali, bensì sarà basata sul *Lemma di Zorn*.

Sul Lemma di Zorn, detto anche Lemma di Kuratowski–Zorn, dai nomi di Kazimierz Kuratowski (1922) e Max Zorn (1935) si basano le dimostrazioni di vari basilari teoremi di esistenza fra i quali, in algebra, il teorema di esistenza di una chiusura algebrica di un campo arbitrario, il teorema di Krull sull’esistenza di ideali massimali in un anello commutativo unitario, ed il teorema di esistenza di una base di Hamel di uno spazio vettoriale; in topologia, il teorema di Tychonoff; in analisi funzionale, il teorema di Hahn–Banach. Come vedremo nel Teorema 7.2.7, anche l’esistenza di ultrafiltri non principali su $\mathbb{N}$ può essere dedotta dal Lemma di Zorn. Approfondiamo questo argomento.

Il Lemma di Zorn afferma che:

Teorema 7.2.6 (Lemma di Zorn) *Se X è un insieme non-vuoto parzialmente ordinato tale che ogni catena ha un maggiorante, allora esiste in X almeno un elemento massimale.¹*

¹ Ricordiamo che, dato un insieme X non-vuoto parzialmente ordinato (cf. [18], p. 25), si chiama *catena* un sottoinsieme non-vuoto di X che sia totalmente ordinato, e che un elemento $m \in X$ è detto *massimale* se non esiste $x \in X$ tale che $x > m$.

DIMOSTRAZIONE. La dimostrazione del Lemma di Zorn segue dall'Assioma di Scelta.^{2,3} Per i dettagli della dimostrazione si veda [18]. Ci limitiamo qui a mostrare intuitivamente come l'Assioma di Scelta entri in gioco nella dimostrazione. Supponiamo che il Lemma di Zorn sia falso. Esiste allora un insieme non-vuoto parzialmente ordinato X dove ogni sottoinsieme totalmente ordinato ha un maggiorante, e per ogni elemento ce ne è uno (strettamente) maggiore. Per ogni sottoinsieme T totalmente ordinato possiamo *scegliere* (tramite l'Assioma di Scelta) un elemento $b(T)$ maggiore di tutti gli elementi di T , in quanto T ha una maggiorazione, e c'è un elemento maggiore di tale maggiorazione. Usando la funzione b , possiamo definire degli elementi $x_0 < x_1 < x_2 < \dots < x_\omega < x_{\omega+1} < x_{\omega+2} < \dots < x_{2\omega} < \dots$ di X . Un tale elenco è veramente “molto lungo”: gli indici non sono numeri naturali, ma *tutti i numeri ordinali*. Esemplifichiamo: partiamo da a_0 arbitrario; allora $\{x_0\}$ è una catena (piuttosto banale), *scegliamo* $x_1 > x_0$ ed aggiungiamolo alla lista; *scegliamo* a_2 maggiore degli elementi di $\{x_0, x_1\}$ ed aggiungiamolo alla lista; ecc.; otteniamo una successione $x_0 < x_1 < x_2 < \dots$ (una vera successione, a indici numeri naturali). Consideriamo poi l'ordinale ω : definiamo

² Si può provare che il Lemma di Zorn, nell'ambito della teoria ZFC degli insiemi, è equivalente all'Assioma di Scelta, nel senso che $ZFC \Rightarrow$ Lemma di Zorn, e se ZF indica ZFC meno AC, allora $ZF +$ Lemma di Zorn $\Rightarrow$ AC.

³ Come è ben noto, l'utilizzo dell'Assioma di Scelta nelle dimostrazioni dei teoremi è stato lungamente posto in discussione. Già O. Zariski, nel 1926, in relazione al teorema che afferma che ogni insieme infinito contiene un insieme numerabile, scriveva ([16], Nota II):

Il punto debole di questo ragionamento è contenuto nel modo stesso in cui è definita la serie estratta. Il susseguirsi dei suoi termini non è determinato da nessuna legge, ma è subordinato invece ad un processo di *infinite scelte arbitrarie*, cioè presuppone un'operazione di un nuovo genere che supera assolutamente le capacità del pensiero umano. Pertanto si può ritenere che la serie estratta si mal definita e che, per il modo con cui è stata costruita, la sua esistenza sia illusoria.

Ancora nel 1948–49, nelle Lezioni tenute alla Scuola Normale Superiore da Carlo Miranda su problemi di esistenza in analisi funzionale, lezioni che ebbero una notevole influenza sulla ricerca in analisi nei decenni seguenti, si legge:

Nel seguito, quando ciò sia indispensabile, noi ci varremo dell'assioma delle infinite scelte, [...] avvertendo, almeno quasi sempre, che la dimostrazione è di carattere zermeliano. Cercheremo però di evitare l'uso dell'assioma quando ciò sia possibile.

Successivamente però, tali scrupoli sono scomparsi.

$x_\omega = b(\{x_n \mid n < \omega\})$, ossia scegliamo x_ω maggiore di tutti i termini della successione (che è una catena), ecc. Alla fine si ha un elenco (dove gli elementi elencati sono tutti diversi) che è “troppo lungo” rispetto agli elementi disponibili, in quanto, intuitivamente, ci sono molti più indici ordinali di quanti possano essere gli elementi di X ; da questa contraddizione segue la tesi. $\square$

7.2.3 Ultrafiltri non principali

Teorema 7.2.7 *Esistono ultrafiltri su $\mathbb{N}$ contenenti il filtro di Fréchet. Tali ultrafiltri non sono principali.*

DIMOSTRAZIONE. Sia $\mathbf{F}$ l'insieme di tutti i filtri su $\mathbb{N}$ che contengono il filtro di Fréchet, ordinato parzialmente per inclusione. Ogni sottoinsieme $\mathbf{C}$ di $\mathbf{F}$ totalmente ordinato è maggiorato da $\bigcup_{C \in \mathbf{C}} C$, che è un filtro che contiene il filtro di Fréchet. Per il Lemma di Zorn (cf. Sezione 7.2.2), $\mathbf{F}$ ha almeno un elemento massimale $\mathcal{U}$.

Se $\mathcal{U}$ fosse principale, e constasse dei sottoinsiemi cui appartiene un dato $m \in \mathbb{N}$, avremmo $\{m\} \in \mathcal{U}$; ma dovendo $\mathcal{U}$ contenere il filtro di Fréchet, dovremmo avere anche $\mathcal{C}\{m\} = \mathbb{N} \setminus \{m\} \in \mathcal{U}$, contro la (F2). $\square$

È anche possibile provare che se $\mathcal{U}$ è un ultrafiltro su $\mathbb{N}$ che non è principale, allora $\mathcal{U}$ contiene il filtro di Fréchet. Quindi le espressioni *ultrafiltro su $\mathbb{N}$ che contiene il filtro di Fréchet* ed *ultrafiltro non principale su $\mathbb{N}$* sono sinonimi.

7.2.4 Gli iperreali

Si è detto che l'idea sottostante alla costruzione dei numeri iperreali è quella di definire “zero” quelle successioni $(p_k)_{k \geq 0}$ di numeri reali che sono 0 “all'infinito”. Naturalmente il problema è quello di formalizzare correttamente il significato dell'espressione “all'infinito”. Data una successione $(p_k)_{k \geq 0}$ di numeri reali, indichiamo con

$$z(p) = \{k \in \mathbb{N} \mid p_k = 0\}$$

l'insieme degli indici $k \in \mathbb{N}$ tali che $p_k = 0$. Procedendo per tentativi ed errori, un primo tentativo potrebbe essere quello di definire “zero” le successioni $(p_k)_{k \geq 0}$ di numeri reali che sono definitivamente uguali a 0, ossia le successioni p per le quali l'insieme di indici $z(p)$ appartiene

al filtro di Fréchet. Tuttavia, in base a questa “definizione”, sarebbero diverse da “zero” le successioni p' e p'' definite da:

$$p'_k = \begin{cases} 0 & \text{se } k \text{ è pari} \\ 1 & \text{se } k \text{ è dispari} \end{cases} \quad p''_k = \begin{cases} 1 & \text{se } k \text{ è pari} \\ 0 & \text{se } k \text{ è dispari} \end{cases}$$

mentre si ha $p'_k p''_k = 0$ per ogni k , ossia $p'p'' = (0)$. Servirà quindi che, in casi come questo, almeno una delle due successioni venga dichiarata “zero”.

Fissiamo allora un ultrafiltro $\mathcal{U}$ su $\mathbb{N}$ che contiene il filtro di Fréchet, e dichiariamo “zero” una successione $(p_k)_{k \geq 0}$ di reali tale che

$$z(p) = \{k \in \mathbb{N} \mid p_k = 0\} \in \mathcal{U}$$

L’idea è di ripetere, *mutatis mutandis*, la costruzione di Méray–Cantor, come descritto sinteticamente nel seguente schema:

Costruzione	Méray–Cantor	iperreali
Campo di partenza	$\mathbb{Q}$	$\mathbb{R}$
Anello di partenza	Successioni di Cauchy in $\mathbb{Q}$	Successioni in $\mathbb{R}$
Ideale massimale I	Successioni $p_n \rightarrow 0$	$p_n = 0$ per $n \in A \in \mathcal{U}$
Quoziente A/I	$\mathbb{R}$	${}^*\mathbb{R}$

Formalizziamo:

Teorema 7.2.8 *Sia I l’insieme delle successioni $p \in \mathbb{R}^{\mathbb{N}}$ tali che*

$$z(p) = \{k \in \mathbb{N} \mid p_k = 0\} \in \mathcal{U}$$

con $\mathcal{U}$ il fissato ultrafiltro su $\mathbb{N}$ che contiene il filtro di Fréchet. Tale I è un ideale massimale dell’anello $A = \mathbb{R}^{\mathbb{N}}$. Il campo ${}^\mathbb{R} = A/I$ è il campo iperreale, o campo dei numeri iperreali.*

DIMOSTRAZIONE. La dimostrazione segue *verbatim* la dimostrazione del Teorema 3.2.5. Verifichiamo che I è un ideale. Siano $p, q \in I$, ossia $z(p), z(q) \in \mathcal{U}$. Ora, per ogni indice k in corrispondenza al quale si ha simultaneamente $p_k = 0$ e $q_k = 0$, è nulla la differenza $p_k - q_k$; pertanto

$z(p) \cap z(q) \subseteq z(p - q)$. Per la (F2) otteniamo che $z(p) \cap z(q) \in \mathcal{U}$, e per la (F1) che $z(p - q) \in \mathcal{U}$, ossia $p - q \in I$. Sia $p \in I$, ossia $z(p) \in \mathcal{U}$, e sia $q \in A$. Ora, per ogni indice k in corrispondenza al quale si ha $p_k = 0$, è nullo il prodotto $p_k q_k$; pertanto $z(p) \subseteq z(pq)$. Per la (F1) otteniamo che $z(pq) \in \mathcal{U}$, ossia $pq \in I$. Quindi I è un ideale. Proviamo che I è un ideale massimale. Sia J un ideale di A , con $I \subseteq J \subseteq A$, e supponiamo sia $I \neq J$. Esiste quindi $w \in J$, ma $w \notin I$, ossia $z(w) \notin \mathcal{U}$. Consideriamo una successione u definita da

$$u_n = \begin{cases} 0 & \text{se } n \in \mathcal{C}z(w) \\ 1 & \text{se } n \in z(w) \end{cases}$$

Dato che $\mathcal{U}$ è un ultrafiltro, si ha $u \in I$, e quindi (per essere $I \subseteq J$) $u \in J$ e quindi la somma $v = u + w \in J$. Sia ora v' la successione ottenuta invertendo termine a termine gli elementi di v (che sono tutti $\neq 0$). Per definizione di ideale, il prodotto termine a termine $v'v$, che è la successione costante $= 1$, appartiene ancora a J e quindi $J = A$. $\square$

Indichiamo con $[p] = p + I$ la classe di equivalenza di p , ossia il numero iperreale rappresentato dalla successione p .⁴

7.2.5 “Unicità”

La costruzione del campo iperreale avviene utilizzando come ingredienti essenziali il campo $\mathbb{R}$ ed un fissato ultrafiltro $\mathcal{U}$ su $\mathbb{N}$ contenente il filtro di Fréchet. Ora, nel mentre $\mathbb{R}$ è “unico” a meno di isomorfismi ordinati (Teorema 2.5.7), per $\mathcal{U}$ sono a priori possibili più scelte.⁵ Si pone pertanto in modo naturale la questione dell’“unicità” (a meno di isomorfismi ordinati) del campo iperreale.

La risposta a questa domanda dipende dall’aggiunta o meno alla teoria degli insiemi adottata, tipicamente la teoria ZFC, di un ulteriore assioma: l’*Ipotesi del Continuo*, indicata con CH. L’Ipotesi del Continuo

⁴ Si noti subito che la $(0, 1, 0, 1, \dots) \cdot (1, 0, 1, 0, \dots) = (0)$ non porta più ad una contraddizione, in quanto essendo che o l’insieme dei numeri pari oppure l’insieme dei numeri dispari deve appartenere a $\mathcal{U}$ (anche se non sappiamo quale dei due appartenga realmente a $\mathcal{U}$), si ha che $[(0, 1, 0, 1, \dots)] = [(0)]$ oppure $[(1, 0, 1, 0, \dots)] = [(0)]$.

⁵ Ad esempio $\mathcal{U}$ potrebbe essere un ultrafiltro (contenente il filtro di Fréchet) cui appartiene l’insieme dei numeri naturali pari, oppure un ultrafiltro (contenente il filtro di Fréchet) cui appartiene l’insieme dei numeri naturali dispari.

afferma che ogni sottoinsieme $E \subseteq \mathbb{R}$ che sia infinito ma non sia numerabile⁶ è necessariamente in corrispondenza biunivoca con $\mathbb{R}$.⁷ È stato dimostrato (cf. [23], p. 33) che se si aggiunge CH agli assiomi della teoria degli insiemi ZFC generalmente accettata, allora la costruzione del campo iperreale è “indipendente” dalla scelta dell’ultrafiltro $\mathcal{U}$, nel senso che nella teoria ZFC+CH due qualunque campi iperreali (costruiti tramite ultrafiltri diversi) sono ordinatamente isomorfi. Se invece non si aggiunge CH a ZFC la risposta alla questione dell’“unicità” (a meno di isomorfismi ordinati) del campo iperreale rimane indeterminata.

7.2.6 Immersione dei reali

Per ogni numero reale $r \in \mathbb{R}$ possiamo costruire una successione costante (r) con ogni termine uguale ad r , e quindi definire un numero iperreale

$$\hat{r} = (r) + I$$

Viene così definita una applicazione $\phi : r \in \mathbb{R} \mapsto \hat{r} \in {}^*\mathbb{R}$ che è iniettiva. Infatti, se in $\mathbb{R}$ si ha $r_1 \neq r_2$, allora la differenza $r_1 - r_2$ fra le successioni costanti uguali risp. ad r_1 ed r_2 è sempre diversa da 0, cioè è uguale a 0 su un insieme di indici vuoto, e siccome $\emptyset \notin \mathcal{U}$, si ha $[r_1] \neq [r_2]$; possiamo quindi *immergere* $\mathbb{R}$ in ${}^*\mathbb{R}$ identificando $\mathbb{R}$ con il sottoinsieme $\hat{\mathbb{R}}$ di ${}^*\mathbb{R}$. In modo evidente ϕ è un omomorfismo, per cui $\phi : \mathbb{R} \rightarrow \hat{\mathbb{R}}$ è un isomorfismo. I numeri iperreali del tipo $\hat{r}$ vengono detti anche *reali standard*, gli altri vengono detti *reali non standard*. Talvolta nel seguito, per non appesantire le notazioni, identificheremo $\hat{\mathbb{R}}$ con $\mathbb{R}$. Se esistano o meno tali reali non standard rimane un punto da chiarire, e sarà chiarito fra poco.

7.2.7 La struttura d’ordine

Per introdurre l’ordinamento nel campo ${}^*\mathbb{R}$ conviene usare, come nel caso del modello di Méray–Cantor, il metodo dell’Osservazione 2.2.3.

⁶ Ricordiamo che gli insiemi *numerabili* sono quelli che possono essere posti in corrispondenza biunivoca con $\mathbb{N}$, ossia quelli per i quali esiste un’applicazione *biiettiva* $\mathbb{N} \rightarrow E$.

⁷ Un insieme E è in corrispondenza biunivoca con $\mathbb{R}$ se esiste un’applicazione *biiettiva* $\mathbb{R} \rightarrow E$.

Definizione 7.2.9 (Definizione dei numeri iperreali positivi) *Sia a un numero iperreale rappresentato dalla successione $(p_k)_{k \geq 0}$. Diciamo che a è positivo, in simboli, $a > 0$ se*

$$z^+(p) = \{k \in \mathbb{N} \mid p_k > 0\} \in \mathcal{U} \quad (7.2)$$

Indichiamo con ${}^*\mathbb{P}$ il sottoinsieme degli elementi positivi di ${}^*\mathbb{R}$.

Si tratta ora di verificare che la definizione non dipende dal rappresentante scelto. Infatti, sia a un numero iperreale rappresentato sia dalla successione $(p_k)_{k \geq 0}$ che dalla $(q_k)_{k \geq 0}$. Per ipotesi, $p - q \in I$, ossia $z(p - q) = \{k \in \mathbb{N} \mid p_k = q_k\} \in \mathcal{U}$. Se $z^+(p) = \{k \in \mathbb{N} \mid p_k > 0\} \in \mathcal{U}$, per la (F2) si ha $z(p - q) \cap z^+(p) \in \mathcal{U}$, ed essendo $z(p - q) \cap z^+(p) \subseteq z^+(q)$, per la (F1) si ha $z^+(q) \in \mathcal{U}$. $\square$

Dimostriamo ora che:

Teorema 7.2.10 *Si ha:*

- (P0) $0 \notin {}^*\mathbb{P}$
- (P1) Se $a, b \in {}^*\mathbb{P}$, allora $a + b, ab \in {}^*\mathbb{P}$.
- (P2) Se $b \neq 0$, vale una ed una sola delle relazioni: $b \in {}^*\mathbb{P}$ oppure $-b \in {}^*\mathbb{P}$.

DIMOSTRAZIONE. (P0) Il numero $0 \in {}^*\mathbb{R}$ è rappresentato da una qualunque successione di numeri reali $(p_k)_{k \geq 0}$ tale che $\{k \in \mathbb{N} \mid p_k = 0\} \in \mathcal{U}$. Se $0 \in {}^*\mathbb{R}$ fosse positivo, avremmo $\{k \in \mathbb{N} \mid p_k > 0\} \in \mathcal{U}$; per la (F1) sarebbe anche $\{k \in \mathbb{N} \mid p_k \neq 0\} \in \mathcal{U}$, una contraddizione con la (F4).

(P1) Siano $a, b > 0$, rappresentati rispettivamente dalle successioni p e q . Da $z^+(p), z^+(q) \in \mathcal{U}$ segue per la (F2) che $\{k \in \mathbb{N} \mid p_k > 0, q_k > 0\} \in \mathcal{U}$, e per la (F1) si ha $z^+(p + q), z^+(pq) \in \mathcal{U}$, ossia $a + b, ab \in {}^*\mathbb{P}$.

(P2) Sia b rappresentato dalla successione p . Poniamo:

$$A = \{k \in \mathbb{N} \mid p_k < 0\}, \quad B = \{k \in \mathbb{N} \mid p_k = 0\}, \quad C = \{k \in \mathbb{N} \mid p_k > 0\}$$

Se $b \neq 0$, si ha $B \notin \mathcal{U}$. Sia $b \notin {}^*\mathbb{P}$, quindi $C \notin \mathcal{U}$. Per la (F4) si ha $CB, CC \in \mathcal{U}$, e per la (F2) risulta $CB \cap CC = C(B \cup C) = A \in \mathcal{U}$, ossia $-b \in {}^*\mathbb{P}$. $\square$

Teorema 7.2.11 *Ponendo*

$$\begin{cases} a < b & \Longleftrightarrow & b - a \in {}^*\mathbb{P} \\ a \leq b & \Longleftrightarrow & ((a < b) \text{ oppure } (a = b)) \end{cases}$$

si ottiene un ordine totale che fa di ${}^*\mathbb{R}$ un campo ordinato. L'isomorfismo $\phi : r \in \mathbb{R} \mapsto \hat{r} \in \hat{\mathbb{R}} \subseteq {}^*\mathbb{R}$ (l'immersione di $\mathbb{R}$ in ${}^*\mathbb{R}$) è un isomorfismo ordinato.

DIMOSTRAZIONE. L'Osservazione 2.2.3 implica che ${}^*\mathbb{R}$ è un campo ordinato. Dati due numeri reali $r < s$, per confrontare gli iperreali $\hat{r}$ e $\hat{s}$ basta osservare che la differenza $\hat{s} - \hat{r}$ è rappresentata dalla successione $(s - r, s - r, s - r, \dots)$ che ha termini positivi in ogni indice $k \in \mathbb{N}$, ed essendo $\mathbb{N} \in \mathcal{U}$ si ha $\hat{s} - \hat{r} > 0$. $\square$

Proveremo ora che ${}^*\mathbb{R}$ non è archimedeo. Conseguentemente ${}^*\mathbb{R}$ ed $\mathbb{R}$ non possono essere ordinatamente isomorfi.⁸

7.2.8 Esistenza degli infinitesimi

Definizione 7.2.12 *Un numero iperreale x è detto infinitesimo se per ogni reale standard $\hat{r} > 0$ si ha $-\hat{r} \leq x \leq \hat{r}$.*

Si noti che, in base alla definizione, 0 è infinitesimo; il teorema seguente assicura che esistono infinitesimi non-nulli:

Teorema 7.2.13 *Esiste un numero iperreale ϵ infinitesimo non-nullo, e quindi il campo ${}^*\mathbb{R}$ non è archimedeo.*

DIMOSTRAZIONE. Sia ϵ il numero iperreale rappresentato dalla successione

$$\left(\frac{1}{k+1} \right)_{k \geq 0} = \left(1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{1}{5}, \frac{1}{6}, \dots \right)$$

Un qualunque reale standard $\hat{r} > 0$ è rappresentato dalla successione costante uguale ad $r > 0$. Si ha certo che $1/(k+1) \leq r$ per tutti gli

⁸ Infatti, se $\mathbb{K}_1$ e $\mathbb{K}_2$ sono campi ordinati, $\mathbb{K}_2$ è archimedeo, ed esiste un isomorfismo ordinato $\phi : \mathbb{K}_1 \rightarrow \mathbb{K}_2$, allora anche $\mathbb{K}_1$ è archimedeo. Dato infatti $a \in \mathbb{K}_1$, $a > 0$, si ha che $\phi(a) > 0$ ed esiste un intero n tale che $n\phi(a) > 1$ in $\mathbb{K}_2$. Applicando l'isomorfismo ordinato ϕ^{-1} si ottiene che $na > 1$ in $\mathbb{K}_1$.

indici $k \geq 1/r$, ossia per tutti gli indici k di un elemento $A \in \mathcal{U}$ (dato che $\mathcal{U}$ contiene il filtro di Fréchet). Lo stesso ragionamento mostra che per ogni n naturale, il multiplo $n\epsilon$ è infinitesimo, e quindi $n\epsilon < 1$, ossia ${}^*\mathbb{R}$ non è archimedeo. $\square$

Indichiamo con $\mathcal{S}$ l'insieme degli infinitesimi (0 compreso).

Naturalmente il teorema precedente è cruciale per affermare che ${}^*\mathbb{R}$ è una estensione ordinata *propria* di $\mathbb{R}$: si noti che a tal fine ci si serve del fatto che l'ultrafiltro $\mathcal{U}$ contiene il filtro di Fréchet. Poiché il campo ${}^*\mathbb{R}$ non è archimedeo, *in ${}^*\mathbb{R}$ non può valere la Proprietà di Dedekind*; di ciò possiamo fornire un esempio esplicito:

Esempio 7.2.14 Sia $A \subseteq {}^*\mathbb{R}$ l'insieme costituito da 0 e dagli infinitesimi positivi, e sia B l'insieme dei reali standard x tali che $0 < x \leq 1$. Allora (A, B) è una coppia di classi separate (ovvio, per definizione di infinitesimo), ma non esiste elemento separatore.

Per provarlo, osserviamo come prima cosa che A non ha massimo: dato un infinitesimo $u \in A$, $u > 0$, si ha che $2u$ è un infinitesimo ed è $2u > u$. Sia ora per assurdo s un elemento separatore di (A, B) . Per definizione di elemento separatore, s è $\leq$ di ogni $x \in B$, e quindi s è $\leq$ di ogni reale standard positivo; inoltre s è $\geq$ di ogni elemento di A , e in particolare di 0, e quindi s è $\geq$ di ogni numero reale standard negativo. Pertanto, per definizione, s è un infinitesimo (positivo), e quindi s è il massimo elemento della classe A , assurdo. $\square$

7.2.9 Confronto fra infinitesimi

Due qualunque infinitesimi possono essere confrontati in quanto ${}^*\mathbb{R}$ è totalmente ordinato, ma con qualche piccola sorpresa. Per esempio, consideriamo ϵ^2 , che è rappresentato dalla successione:

$$1, \frac{1}{2^2}, \frac{1}{3^2}, \frac{1}{4^2}, \frac{1}{5^2}, \frac{1}{6^2}, \frac{1}{7^2}, \frac{1}{8^2}, \frac{1}{9^2}, \frac{1}{10^2}, \frac{1}{11^2}, \dots$$

È banalmente vero che $\epsilon^2 < \epsilon$, ma se nella successione $(\frac{1}{k+1})_{k \geq 0}$ che rappresenta ϵ eleviamo al cubo i termini di indice dispari, il risultato:

$$1, \frac{1}{2^3}, \frac{1}{3}, \frac{1}{4^3}, \frac{1}{5}, \frac{1}{6^3}, \frac{1}{7}, \frac{1}{8^3}, \frac{1}{9}, \frac{1}{10^3}, \frac{1}{11}, \dots$$

rappresenta un infinitesimo η del quale non sappiamo dire se è $\eta < \epsilon^2$ oppure $\eta > \epsilon^2$, pur essendo vera sicuramente una (e una sola) delle due, in quanto sappiamo che dei due sottoinsiemi di $\mathbb{N}$, quello dei numeri pari e quello dei numeri dispari, uno ed uno solo appartiene all'ultrafiltro $\mathcal{U}$, ma non sappiamo quale.

7.2.10 Esistenza degli infiniti

Un numero iperreale x è detto *finito* se esiste un reale standard $\hat{r}$ tale che $-\hat{r} \leq x \leq \hat{r}$; altrimenti, x è detto *infinito*. Ovviamente i reali standard sono finiti. Ma esistono numeri iperreali non-finiti, detti numeri iperreali infiniti: per esempio il numero H rappresentato dalla successione $(1, 2, 3, 4, 5, \dots)$ è un iperreale infinito. Un numero iperreale $H \neq 0$ è infinito se e solo se $\epsilon = 1/H$ è infinitesimo. Lasciamo al lettore scrupoloso le relative verifiche.

7.2.11 Teorema della parte standard

Questo importante teorema afferma sostanzialmente che ogni numero reale *finito* non standard x è “infinitamente vicino” (nel senso che la differenza è infinitesima) ad un numero reale standard $y = \text{st}(x)$, che viene detto la *parte standard* dell'iperreale finito x . Riassumendo:

Teorema 7.2.15 (Teorema della parte standard) *I numeri iperreali finiti costituiscono un anello commutativo ed unitario $\mathcal{M}$. L'insieme $\mathcal{S}$ degli infinitesimi è un ideale massimale di $\mathcal{M}$. Dato un numero iperreale finito x , esiste un unico numero reale $y = \text{st}(x)$, detto la parte standard di x , la cui differenza con x è in $\mathcal{S}$, ossia è infinitesima. L'applicazione $\text{st}: \mathcal{M} \rightarrow \mathbb{R}$ è un omomorfismo il cui nucleo è l'ideale $\mathcal{S}$ degli infinitesimi. Il quoziente $\mathcal{M}/\mathcal{S}$ è quindi isomorfo ad $\mathbb{R}$.*

DIMOSTRAZIONE. Per provare che i numeri iperreali finiti costituiscono in ${}^*\mathbb{R}$ un anello $\mathcal{M}$ commutativo ed unitario, basta provare che la differenza ed il prodotto di numeri iperreali finiti sono numeri iperreali finiti: lasciamo al lettore scrupoloso la verifica di questa affermazione (come di altre del tutto simili in seguito).

Proviamo ora che l'insieme $\mathcal{S}$ degli infinitesimi è un ideale massimale dell'anello $\mathcal{M}$. Infatti, $\mathcal{S}$ è un ideale, in quanto la differenza di due infinitesimi è infinitesima, ed il prodotto di un infinitesimo per un qualunque numero *finito* è infinitesimo. Infine, se J è un ideale di $\mathcal{M}$ che contiene propriamente $\mathcal{S}$, esiste un elemento $x \in J$ non infinitesimo: tale x non è zero, né infinitesimo, per cui esiste il suo inverso $1/x$ e tale inverso non è infinito. Quindi $1 \in J$, e l'ideale $\mathcal{S}$ è massimale.

Il quoziente $\mathcal{M}/\mathcal{S}$ ripartisce gli iperreali finiti in classi di equivalenza: due iperreali nella stessa classe differiscono per un infinitesimo. Mostriamo che ogni classe di equivalenza $\langle x \rangle$ contiene uno ed un solo numero reale standard, che indicheremo con $y = \text{st}(x)$, e che diremo *la parte standard* di x (e di qualunque altro iperreale finito della classe considerata).

In $\langle x \rangle$ non possono stare due diversi reali standard, in quanto la differenza fra due reali standard diversi non è infinitesima. Sia quindi $\langle x \rangle = x + \mathcal{S}$ una delle classi di equivalenza.⁹ Se x è un reale standard poniamo $y = \text{st}(x) = x$. Altrimenti esistono reali standard sia minori che maggiori di x . Ripartiamo quindi i reali standard in due classi: $A = \{a \in \mathbb{R} \mid a < x\}$ e $B = \{b \in \mathbb{R} \mid b > x\}$, che risultano separate. Per la completezza alla Dedekind di $\mathbb{R}$, o A ha massimo oppure B ha minimo. Supponiamo per esempio $y = \min(B)$. Per ogni $a \in A$, $b \in B$, si ha $a < x < b$, quindi $a - y < x - y < b - y$. La differenza $x - y$ è un infinitesimo in quanto $a - y$ è un generico reale negativo e $b - y$ è un generico reale positivo. Pertanto y è un reale standard appartenente alla classe $\langle x \rangle$, e possiamo definire $y = \text{st}(x)$.

Proviamo che $\text{st}(x + y) = \text{st}(x) + \text{st}(y)$. Consideriamo le classi $\langle x \rangle$, $\langle y \rangle$ e la loro somma $\langle x \rangle + \langle y \rangle = \langle x + y \rangle$. Ora $\text{st}(x) \in \langle x \rangle$, $\text{st}(y) \in \langle y \rangle$, quindi, per costruzione $\text{st}(x) + \text{st}(y) \in \langle x + y \rangle$; ma $\text{st}(x) + \text{st}(y)$ è standard, e l'unico reale standard in $\langle x + y \rangle$ è $\text{st}(x + y)$. Analogamente si prova che $\text{st}(xy) = \text{st}(x) \cdot \text{st}(y)$. Le restanti affermazioni seguono dal Teorema fondamentale di omomorfismo per gli anelli (cf. [18], p. 221). $\square$

In base al Teorema della parte standard, ogni numero iperreale finito x della classe di equivalenza $\langle x \rangle$ si “decompone” nella somma della *parte standard* $\text{st}(x)$ e della *parte infinitesima* $\sigma \in \mathcal{S}$:

$$x = \text{st}(x) + \sigma$$

⁹ Utilizziamo un simbolo $\langle \rangle$ per le classi dell'equivalenza modulo $\mathcal{S}$ diverso da quello $[]$ usato per le classi dell'equivalenza modulo I .

Al variare dell'infinitesimo σ , i numeri $\text{st}(x) + \sigma$ costituiscono un insieme $\mathcal{X}$ detto *monade* o *nuvola* del numero reale standard $\text{st}(x)$, e che altro non è che una classe dell'equivalenza modulo $\mathcal{S}$. Naturalmente tutti i punti della stessa monade hanno la medesima parte standard.

È interessante notare che la monade $\mathcal{X}$ di un numero reale standard x , pur essendo limitata,¹⁰ non ha né estremo inferiore né estremo superiore. Supponiamo per assurdo che esista l'estremo superiore $\lambda = \sup \mathcal{X}$. Dato un infinitesimo $\varepsilon > 0$, esiste un elemento $x' \in \mathcal{X}$ tale che

$$x \leq \lambda - \varepsilon < x' \leq \lambda$$

Ora $0 \leq \lambda - \varepsilon - x < x' - x$, e $x' - x$ è infinitesimo: conseguentemente $\lambda - \varepsilon - x$ e quindi anche $\sigma = \lambda - x$ sono infinitesimi, ossia $\lambda = x + \sigma \in \mathcal{X}$. Pertanto $\lambda = \max \mathcal{X}$, il che è impossibile, in quanto $\lambda < \lambda + \sigma = x + 2\sigma \in \mathcal{X}$. Analogamente si prova che non esiste $\inf \mathcal{X}$. $\square$

7.3 Cenni allo sviluppo dell'analisi non-standard

L'analisi non-standard consente in taluni casi di studiare funzioni $f : \mathbb{R} \rightarrow \mathbb{R}$ “lavorando” nell'estensione ${}^*\mathbb{R}$ di $\mathbb{R}$ e poi “riportando” il risultato in $\mathbb{R}$.¹¹

Data una funzione $f : \mathbb{R} \rightarrow \mathbb{R}$ possiamo definire un'estensione ${}^*f : {}^*\mathbb{R} \rightarrow {}^*\mathbb{R}$ della funzione f al modo seguente. Dato un qualunque numero iperreale $[x]$ rappresentato da una successione $(x_k)_{k \geq 0}$ di numeri reali, poniamo ${}^*f([x]) = [y]$ dove la successione y è definita per $k \geq 0$ da $y_k = f(x_k)$. Osserviamo che se $(x'_k)_{k \geq 0}$ è un'altra successione di numeri reali che rappresenta lo stesso numero iperreale $[x]$ si ha

$$\{k \in \mathbb{N} \mid x_k = x'_k\} \in \mathcal{U}$$

e poiché

$$\{k \in \mathbb{N} \mid x_k = x'_k\} \subseteq \{k \in \mathbb{N} \mid f(x_k) = f(x'_k)\}$$

¹⁰ Ad esempio, $\mathcal{X}$ è limitata inferiormente da $x - 1$ e superiormente da $x + 1$.

¹¹ Si tratta di una tecnica che non è per nulla originale: ad esempio sappiamo bene che per trovare radici reali di un polinomio reale di terzo grado conviene lavorare in $\mathbb{C}$ e poi riportare il risultato in $\mathbb{R}$.

si ha che

$$\{k \in \mathbb{N} \mid f(x_k) = f(x'_k)\} \in \mathcal{U}$$

ossia il valore $*f([x]) = [y]$ non dipende dal rappresentante utilizzato per $[x]$. Pertanto $*f$ è “ben definita”. Si noti anche che se $\widehat{r}$ è un reale standard rappresentato dalla successione costante uguale ad r , si ha

$$*f(\widehat{r}) = [f(r)] = \widehat{f(r)}$$

per cui l'uso del termine *estensione* è corretto; pertanto, quando non siano possibili fraintendimenti, la funzione estesa $*f$ può venir indicata semplicemente con il solo simbolo f , omettendo l'asterisco.

Mostriamo, come esempio, la definizione non-standard di derivata ed alcune sue immediate conseguenze.

Definizione di derivata

Definizione 7.3.1 Si dice che un numero reale $m = f'(x)$ è la derivata di $f : \mathbb{R} \rightarrow \mathbb{R}$ in $x \in \mathbb{R}$ se per ogni infinitesimo non-nullo Δx si ha

$$\text{st} \left(\frac{f(x + \Delta x) - f(x)}{\Delta x} \right) = m$$

In altri termini, f è derivabile in x se, al variare dell'infinitesimo non-nullo Δx , i vari rapporti incrementali $(f(x + \Delta x) - f(x))/\Delta x$ hanno tutti la medesima parte standard, ovvero appartengono tutti alla stessa monade.

Esempio 7.3.2 Calcoliamo ad esempio la derivata di $f(x) = x^3$. Sia dato un infinitesimo $\Delta x \neq 0$. Calcoliamo:

$$\begin{aligned} \frac{f(x + \Delta x) - f(x)}{\Delta x} &= \frac{(x + \Delta x)^3 - x^3}{\Delta x} = \frac{3x^2\Delta x + 3x\Delta x^2 + \Delta x^3}{\Delta x} \\ &= 3x^2 + 3x\Delta x + \Delta x^2 \end{aligned}$$

Ma

$$\text{st}(3x^2 + 3x\Delta x + \Delta x^2) = 3x^2$$

Quindi $f'(x) = 3x^2$. □

Ancora, e per ultimo, diamo un “assaggio” di qualche sviluppo del calcolo differenziale non-standard. Per ulteriori sviluppi e dettagli, nonché per un itinerario didattico completo, cf. [7], [23], [30], [31].

Formula di approssimazione lineare Supponiamo che $f : \mathbb{R} \rightarrow \mathbb{R}$ sia *derivabile* in x_0 secondo la Definizione 7.3.1. Sia $\mathcal{X}$ la monade di x_0 e sia $x \in \mathcal{X}$. Se $x \neq x_0$ la differenza $\Delta x = x - x_0$ è un infinitesimo non-nullo. Poniamo:

$$\Delta y = f(x) - f(x_0)$$

Visto che, per la Definizione 7.3.1,

$$f'(x_0) = \text{st} \left(\frac{\Delta y}{\Delta x} \right)$$

esiste un infinitesimo σ tale che

$$\frac{\Delta y}{\Delta x} = f'(x_0) + \sigma$$

Queste relazioni possono venir riassunte nel:

Teorema 7.3.3 (Formula di approssimazione lineare) *Sia $f : \mathbb{R} \rightarrow \mathbb{R}$ derivabile in x_0 . Per ogni x nella monade $\mathcal{X}$ di x_0 esiste un infinitesimo $\sigma = \sigma(x)$ (dipendente da x) tale che*

$$f(x) = f(x_0) + f'(x_0)(x - x_0) + \sigma(x)(x - x_0) \quad (7.3)$$

Si può definire su $\mathcal{X}$ una funzione “affine” r ponendo, per $x \in \mathcal{X}$ (x iperreale)

$$r(x) = f(x_0) + f'(x_0)(x - x_0)$$

dove $\Delta x = x - x_0$ è un infinitesimo. Per $x \in \mathcal{X}$ il valore $f(x)$ differisce da $r(x)$ per il termine $\sigma(x)(x - x_0)$, che, diviso per $(x - x_0)$ (quando $x \neq x_0$), dà per quoziente l’infinitesimo $\sigma(x)$. La funzione r è l’*approssimante lineare* di f in x_0 . L’approccio non-standard pone quindi in particolare evidenza il carattere *locale* dell’approssimazione lineare, dato che r è definita in modo “naturale” per i soli x *infinitamente vicini* ad x_0 , ossia quelli appartenenti alla sua monade.¹²

¹² Nel caso standard, la formula di approssimazione lineare per una funzione $f : \mathbb{R} \rightarrow \mathbb{R}$ derivabile in x_0 si scrive in modo del tutto analogo alla (7.3), con la fondamentale differenza che σ non è un numero iperreale infinitesimo, ma un numero reale $\sigma = \sigma(x)$

La derivabilità implica la continuità

Definizione 7.3.4 Si dice che una funzione $f : \mathbb{R} \rightarrow \mathbb{R}$ è continua in x_0 se, per ogni x infinitamente vicino a x_0 , si ha che $f(x)$ è infinitamente vicino a $f(x_0)$; equivalentemente, $f : \mathbb{R} \rightarrow \mathbb{R}$ è continua in x_0 se per ogni infinitesimo Δx la differenza $\Delta y = f(x) - f(x_0)$ è un infinitesimo.

Teorema 7.3.5 Una funzione $f : \mathbb{R} \rightarrow \mathbb{R}$ derivabile in x_0 è continua in x_0 .

DIMOSTRAZIONE. La tesi segue immediatamente dalla (7.3), in quanto se $\Delta x = x - x_0$ è un infinitesimo anche $\Delta y = f(x) - f(x_0) = (f'(x_0) + \sigma(x)) \cdot \Delta x$, che è il prodotto di un iperreale finito per un infinitesimo, è un infinitesimo.

che dipende da x , tale che

$$\frac{\sigma(x)}{x - x_0} \rightarrow 0 \text{ per } x \rightarrow x_0$$

Nell'approccio standard l'approssimante lineare $r(x) = f(x_0) + f'(x - x_0)$ è quindi definibile in modo “naturale” per ogni $x \in \mathbb{R}$, ma la differenza $f(x) - r(x)$ può assumere valori “incontrollati” per $x \neq x_0$: di tale differenza sappiamo solo che, se viene divisa per $(x - x_0)$ (quando $x \neq x_0$), dà un quoziente che tende a 0 per $x \rightarrow x_0$.

Capitolo 8

Numeri reali e . . .

*And therefore as a stranger give it welcome.
There are more things in heaven and earth, Horatio,
Than are dreamt of in your philosophy.^a*

SHAKESPEARE,
Hamlet

^a E allora dagli il benvenuto, come si fa con gli stranieri. Vi sono più cose in cielo e in terra, Orazio, di quante non ne sogni la tua filosofia.

8.1 Introduzione

Questo capitolo contiene vari spunti di riflessione sul concetto di numero reale. Gli argomenti presentati mostrano la grande delicatezza, anche in termini filosofici, della nozione di numero reale. Sono per prime esaminate alcune questioni legate alla *non-numerabilità* di $\mathbb{R}$. Vengono richiamati i sottoinsiemi numerabili “classici”, quello dei numeri razionali e quello dei numeri algebrici, e viene introdotto, in modo informale, il sottoinsieme numerabile dei numeri calcolabili. Si mostra che i numeri dei corrispondenti insiemi complementari (irrazionali, trascendenti, non-calcolabili) costituiscono, da diversi punti di vista, la “stragrande maggioranza” dei numeri reali, con la sconcertante conclusione che “quasi tutti” i numeri reali non possono essere né mai saranno “calcolati”

dall'uomo. Questa discussione sui numeri calcolabili e non-calcolabili costituisce un importante argomento logico-informatico, che spesso è assente dalle trattazioni più strettamente matematiche sui numeri reali. Viene anche proposta una “soluzione” del Paradosso di Richard (legato ai numeri “definibili” e “non-definibili”) ispirata direttamente al Teorema di Indecidibilità di Turing. Questa parte può anche essere di stimolo per attività interdisciplinari (con l'informatica, con la logica, con la filosofia). Il capitolo è infine concluso da un contributo di Michela Maschietto relativo alla genesi del concetto di numero reale secondo alcuni recenti studi che utilizzano ed applicano le tecniche delle scienze cognitive e delle neuroscienze. Dal punto di vista dello sviluppo concettuale, i numeri reali hanno una struttura assolutamente diversa da quella della numerosità, legata al contare ed ai numeri naturali, in quanto non avrebbero, a differenza dei numeri naturali, un analogo “neuronale” nel cervello umano. Per la comprensione dei numeri reali vi sarebbe quindi, in aggiunta agli ostacoli *epistemologici* noti, un ostacolo *biologico* non eliminabile.

Negli spunti di questo capitolo non pretendiamo di dare al lettore certezze, ma speriamo perlomeno di generare il “sospetto” che questi temi non siano per nulla banali e scontati.

8.2 Numeri reali e cardinalità

Ricordiamo che gli insiemi *numerabili* sono quelli che possono essere posti in corrispondenza biunivoca con $\mathbb{N}$, ossia quelli i cui elementi possono essere elencati tutti (nessuno escluso e senza ripetizioni) come termini di una successione $(x_i)_{i \geq 0}$. Una celebre argomentazione di Cantor, detta *procedimento diagonale* (vedi Sezione 8.2.3) mostra che $\mathbb{R}$ non è numerabile: intuitivamente, $\mathbb{R}$ ha “molti più elementi” di quanti possano essere elencati in una successione. Si dice che gli insiemi numerabili hanno la *cardinalità* (o la *potenza*) *del numerabile*, e che $\mathbb{R}$ ha la *cardinalità* (o la *potenza*) *del continuo*.¹ Gli spunti intendono evidenziare e far riflettere sul fatto che la non-numerabilità di $\mathbb{R}$ è ben lontana dall'essere una semplice “curiosità matematica”, e conduce a importanti problematiche sia logiche che filosofiche. Non a caso infatti, il procedimento diagonale di Cantor, che in una sua forma elementare dimostra la non-numerabilità

¹ Taluni usano il termine *numerosità* come sinonimo di cardinalità, o potenza.

di $\mathbb{R}$, è, in contesti diversi, la chiave per alcuni profondi teoremi di logica, come il Teorema di Indecidibilità di Turing e lo stesso Teorema di Indecidibilità di Gödel, che qui ci limitiamo a nominare, rimandando a testi di logica matematica per eventuali approfondimenti. Parafrasando Shakespeare, le costruzioni di $\mathbb{R}$ producono alla fine veramente “molte più cose” di quante se ne possano immaginare.² A questo proposito, ci appare molto significativa ed emblematica la seguente citazione di Émile Borel [9], che esprime ancor oggi tutta la sua problematicità:

Molti analisti, in effetti, mettono al primo posto la nozione del continuo, che è quella che interviene, in modo più o meno esplicito, nei loro ragionamenti. Ho indicato di recente^a in cosa tale nozione di continuo, considerata come una potenza superiore a quella del numerabile, mi pare essere una nozione puramente negativa, dato che la potenza degli insiemi numerabili è la sola che possa essere conosciuta in modo positivo, la sola che interviene *effettivamente* nei nostri ragionamenti. È chiaro infatti, che l'insieme degli elementi analitici suscettibili di essere realmente definiti e considerati non può che essere un insieme numerabile; ...³

^a BOREL, *Les paradoxes de la théorie des ensembles*, Annales scientifiques de l'École Normale supérieure, 3^e série, tome XXV (1908).

Riteneva quindi Borel che le cardinalità del numerabile e del continuo non possano esser messe sullo stesso piano, come non possiamo mettere sullo stesso piano la derivabilità e la non-derivabilità, o l'integrabilità e la non-integrabilità di una funzione, e che la cardinalità del continuo sia un fatto *negativo*, visto appunto come mancanza di numerabilità.

² Osserviamo incidentalmente che, se le affermazioni della nostra filosofia sono frasi di un linguaggio naturale, come l'italiano o l'inglese, esse sono necessariamente in quantità al più numerabile.

³ Beaucoup d'analystes, en effet, mettent au premier rang la notion du continu; c'est celle qui intervient d'une manière plus ou moins explicite dans leurs raisonnements. J'ai indiqué récemment [...] en quoi cette notion du continu, considéré comme ayant une puissance supérieure à celle du dénombrable, me paraît être une notion purement négative, la puissance des ensembles dénombrables étant la seule qui nous soit connue d'une manière positive, la seule qui intervienne *effectivement* dans nos raisonnements. Il est clair, en effet, que l'ensemble des éléments analytiques susceptibles d'être réellement définis et considérés ne peut être qu'un ensemble dénombrable; ...

Rimandiamo a [12] per ulteriori aspetti riguardanti il legame fra il concetto di numero reale e la teoria della complessità.

8.2.1 Insiemi numerabili

Ricordiamo che un insieme E è *numerabile* se esiste una applicazione biiettiva $\nu : \mathbb{N} \rightarrow E$ (che viene detta una *numerazione* di E). In altri termini, E è *numerabile* se e solo se esiste una successione $(x_i)_{i \geq 0}$ i cui termini sono tutti distinti e sono tutti e soli gli elementi di E (basta infatti porre $x_i = \nu(i)$).

Proposizione 8.2.1 *La unione numerabile $E = \bigcup_{i \geq 0} E_i$ di insiemi finiti E_i è finita o numerabile.*

DIMOSTRAZIONE. Non è restrittivo supporre che gli insiemi E_i siano tutti non-vuoti. Quindi per ogni $i \geq 0$ esiste una $(n_i + 1)$ -pla $(x_{i0}, x_{i1}, \dots, x_{in_i})$ i cui termini sono tutti e soli gli elementi di E_i . Scrivendo allora

$$x_{00}, x_{01}, \dots, x_{0n_0}, x_{10}, x_{11}, \dots, x_{1n_1}, \dots, x_{i0}, x_{i1}, \dots, x_{in_i}, \dots$$

otteniamo una numerazione di E . □

Dalla Proposizione 8.2.1 segue subito che il prodotto cartesiano $\mathbb{Z} \times \mathbb{Z}$ è numerabile, in quanto

$$\mathbb{Z} \times \mathbb{Z} = \bigcup_{i \geq 0} \{(n_1, n_2) \mid \max\{|n_1|, |n_2|\} = i\}$$

e che, più in generale, per $k \geq 2$, è numerabile il prodotto

$$\mathbb{Z}^k = \underbrace{\mathbb{Z} \times \mathbb{Z} \times \dots \times \mathbb{Z}}_{k \text{ fattori}} = \bigcup_{i \geq 0} \left\{ (n_1, n_2, \dots, n_k) \mid \max_{1 \leq j \leq k} |n_j| = i \right\}$$

Dalla numerabilità di $\mathbb{Z} \times \mathbb{Z}$ segue che

Proposizione 8.2.2 *La unione numerabile $E = \bigcup_{i \geq 0} E_i$ di insiemi numerabili E_i è numerabile.*

DIMOSTRAZIONE. Gli elementi dell'insieme numerabile E_i possono essere elencati in una successione $(x_{ij})_{j \geq 0}$ e quindi quelli della unione $E = \bigcup_{i \geq 0} E_i$ possono essere elencati con un indice $(i, j) \in \mathbb{N} \times \mathbb{N}$, dove il prodotto $\mathbb{N} \times \mathbb{N}$, in quanto sottoinsieme di $\mathbb{Z} \times \mathbb{Z}$, è numerabile. □

8.2.2 Il procedimento diagonale di Cantor

Il *procedimento diagonale di Cantor* è un semplice metodo che consente di costruire (e quindi, in particolare, di dimostrare l'esistenza di) un allineamento binario che *non* appartiene ad un prefissato insieme *numerabile* E di allineamenti binari. Sia dato infatti un insieme numerabile E di allineamenti binari, e sia $\nu : \mathbb{N} \rightarrow E$ una sua numerazione. Possiamo allora considerare una *matrice doppiamente infinita* $A = (a_{ij})_{i \geq 0, j \geq 0}$ di due simboli, zero '0' e uno '1', la cui riga i -esima è l'allineamento i -esimo nella numerazione di E , ossia

$$a_{ij} = (a_{i0}, a_{i1}, a_{i2}, a_{i3}, a_{i4}, a_{i5}, a_{i6}, \dots)$$

Definendo, per ogni $i \geq 0$,

$$\alpha_i = 1 - a_{ii} = \begin{cases} 1 & \text{se } a_{ii} = 0 \\ 0 & \text{se } a_{ii} = 1 \end{cases}$$

di ha un allineamento α che non può coincidere con alcuna riga della matrice A : infatti, quale che sia l'indice i , l'allineamento α non può coincidere con la riga i -esima, avendone diverso, per costruzione, proprio il termine i -esimo. $\square$

8.2.3 Non-numerabilità di $[0, 1)$

L'intervallo $[0, 1) \subseteq \mathbb{R}$ (e di conseguenza $\mathbb{R}$) non è numerabile. In questa sezione presentiamo, in due varianti, una dimostrazione basata sul procedimento diagonale di Cantor.⁴ Le varianti differiscono fra loro per la interpretazione data ad un allineamento binario. Infatti un allineamento binario come ad esempio

$$a = (a_i)_{i \geq 0} = 110010010000111111011010101000100010000101101000 \dots$$

può essere considerato (a) la mantissa di un numero reale $x \in [0, 1)$ scritto in base 2:

$$\begin{aligned} x &= 0.110010010000111111011010101000100010000101101000 \dots_2 \\ &= 0.785398163397 \dots \end{aligned}$$

⁴ Una ulteriore dimostrazione sarà presentata nella Sezione 8.3.1.

oppure (b) la funzione caratteristica di un sottoinsieme $X \subseteq \mathbb{N}$, che, per l'allineamento considerato, è l'insieme:

$$\begin{aligned} X &= \{i \in \mathbb{N} \mid a_i = 1\} \\ &= \{0, 1, 4, 7, 12, 13, 14, 15, 16, 17, 19, 20, 22, 24, 26, 30, 34, 39, \dots\} \end{aligned}$$

Proposizione 8.2.3 *L'intervallo $[0, 1) \subseteq \mathbb{R}$ non è numerabile.*

DIMOSTRAZIONE. (a) Supponiamo per assurdo che esista una numerazione $(x_i)_{i \geq 0}$ di tutti i numeri reali dell'intervallo $[0, 1)$. Possiamo allora considerare una matrice doppiamente infinita le cui righe sono gli allineamenti decimali che costituiscono le mantisse dei numeri reali x_i in base 2. Per ciascuno di quei numeri reali che hanno una doppia rappresentazione (0-periodica e 1-periodica), scriviamo due righe, una per ciascuna delle due rappresentazioni. L'argomento diagonale di Cantor costruisce un allineamento decimale che rappresenta un numero reale dell'intervallo $[0, 1)$ che non è presente nella lista, il che è una contraddizione. $\square$

La seconda versione della dimostrazione è più “impegnativa” dal punto di vista logico.

DIMOSTRAZIONE. (b) Supponiamo per assurdo, come nella dimostrazione precedente, di poter numerare *tutti* gli allineamenti binari. Dato il generico allineamento i -esimo della lista, consideriamo il sottoinsieme $X_i \subseteq \mathbb{N}$ che ha l'allineamento i -esimo dato come funzione caratteristica. Definiamo ora un sottoinsieme $N \subseteq \mathbb{N}$ ponendo

$$N = \{i \mid i \notin X_i\}$$

La funzione caratteristica di questo sottoinsieme N deve allora essere presente nella lista, per esempio nella riga n -esima, ossia deve essere $N = X_n$ per un $n \in \mathbb{N}$ opportuno. In base alle definizioni date, si ha $n \in N \iff n \notin N$. Abbiamo quindi ottenuto un “insieme assurdo”, ed il teorema è provato. $\square$

Osservazione 8.2.4 La discussione precedente mostra che il procedimento diagonale di Cantor contiene una argomentazione *impredicativa* (cf. [33], pag. 1184); per tale motivo è guardato “con cautela” da taluni filosofi della matematica. Notiamo ancora che la definizione dell’“insieme assurdo” N ricorda

la definizione dell’“insieme assurdo” R del Paradosso di Russel.⁵ Intuitivamente, R è l’insieme di tutti gli insiemi che non sono elementi di se stessi. Formalmente, $R = \{x \mid x \notin x\}$. La considerazione di R porta all’assurdo che $R \in R \iff R \notin R$.

8.2.4 Numeri razionali, numeri algebrici, numeri calcolabili

In questa sezione richiamiamo alcuni importanti sottoinsiemi numerabili di $\mathbb{R}$. Si tratta di:

- Il sottoinsieme $\mathbb{Q}$ dei numeri *razionali*.
- Il sottoinsieme $\mathcal{A}$ dei numeri *algebrici*. Ricordiamo che un numero reale x è *algebrico* se esiste un polinomio p a coefficienti interi tale che $p(x) = 0$.
- Il sottoinsieme $\mathcal{T}$ dei numeri *calcolabili*. La definizione di numero calcolabile è dovuta ad Alan Turing [57]. Intuitivamente, un numero reale x , per esempio appartenente all’intervallo $[0, 1)$, è *calcolabile* se esiste un algoritmo, sintatticamente corretto, che prende come input un intero $n \geq 1$ e termina la sua esecuzione fornendo come output le prime n cifre di una rappresentazione binaria o decimale di x .

Proposizione 8.2.5 *I sottoinsiemi $\mathbb{Q}$ dei numeri razionali, $\mathcal{A}$ dei numeri algebrici, e $\mathcal{T}$ dei numeri calcolabili sono sottoinsiemi numerabili.*

DIMOSTRAZIONE. (a) Rappresentiamo ogni $r \in \mathbb{Q}$ con una frazione $r = p/q$ ridotta ai minimi termini, con $q > 0$; facendo corrispondere ad r la coppia $(p, q) \in \mathbb{Z} \times \mathbb{Z}$, possiamo pensare $\mathbb{Q}$ come sottoinsieme dell’insieme numerabile $\mathbb{Z} \times \mathbb{Z}$, e pertanto $\mathbb{Q}$ è numerabile.

(b) L’insieme P_n dei polinomi a coefficienti interi di grado $\leq n$ è numerabile, in quanto un polinomio $p = a_0 + a_1x + \dots + a_nx^n$ a coefficienti interi di grado $\leq n$ è identificabile con la $(n+1)$ -pla $(a_0, a_1, \dots, a_n) \in \mathbb{Z}^{n+1}$ dei suoi coefficienti, e tali $(n+1)$ -ple costituiscono un insieme numerabile. Inoltre, poiché ogni polinomio $p \in P_n$ ha al più un numero

⁵ Bertrand Russell, 1901.

finito di radici, possiamo rappresentare l'insieme $\mathcal{A}$ dei numeri algebrici come

$$\mathcal{A} = \bigcup_{n \geq 0} \bigcup_{p \in P_n} \{x \in \mathbb{R} \mid p(x) = 0\}$$

ossia come unione numerabile di insiemi finiti. Per la Proposizione 8.2.2, si ha che $\mathcal{A}$ è numerabile.

(c) Un algoritmo che possa calcolare le cifre binarie di un numero calcolabile consiste in una stringa finita di codice scritto con un alfabeto finito. Se l'alfabeto ha, ad esempio, due soli simboli, 0 ed 1, gli algoritmi scritti con esattamente N simboli sono al massimo 2^N . L'insieme di tutti gli algoritmi in questione è dunque unione numerabile di insiemi finiti. Per la Proposizione 8.2.2, si ha che $\mathcal{T}$ è numerabile. $\square$

8.2.5 Numeri irrazionali, numeri trascendenti, numeri non-calcolabili

I numeri reali che non sono razionali si dicono, come si sa, *irrazionali*; i numeri reali che non sono algebrici si dicono *trascendenti*; i numeri reali che non sono calcolabili si dicono *non-calcolabili*. L'*esistenza* di questi tre tipologie di numeri (irrazionali, trascendenti, non-calcolabili) discende immediatamente dalla non-numerabilità di $\mathbb{R}$: essendo che gli insiemi dei razionali, degli algebrici, e dei calcolabili sono insiemi numerabili, devono necessariamente esistere numeri reali che non sono razionali, devono necessariamente esistere numeri reali che non sono algebrici, e devono necessariamente esistere numeri reali che non sono calcolabili, altrimenti $\mathbb{R}$ sarebbe numerabile.

La esplicita esibizione di un numero irrazionale è semplice, per esempio $\sqrt{2}$ è irrazionale (cf. Capitolo 1). È invece difficile decidere se un dato numero reale x sia algebrico oppure trascendente (cf. [33], p. 980). Si sa ad esempio che il numero di Nepero e è trascendente (Hermite [25], 1873) e che π è trascendente (Lindemann [36], 1882). Ancora più difficile è esibire esplicitamente un numero non-calcolabile (malgrado il fatto che, come vedremo, essi costituiscano la “stragrande maggioranza” dei numeri reali).

8.2.6 Un esempio di numero reale non-calcolabile

Per esibire esplicitamente un numero non-calcolabile può venire l'idea di utilizzare direttamente il procedimento diagonale di Cantor. Dato che i numeri calcolabili dell'intervallo $[0, 1)$ sono numerabili, possiamo applicare l'argomento diagonale di Cantor all'elenco delle loro rappresentazioni binarie, ottenendo in tal modo la rappresentazione binaria di un numero r dell'intervallo $[0, 1)$ che non è calcolabile. Questa affermazione può essere “inquietante”: da un lato si afferma che l'argomento diagonale definisce le cifre di r , dall'altro si afferma che r non è calcolabile (e sarebbe quindi un numero assurdo, simultaneamente calcolabile e non-calcolabile). Intuitivamente, l'errore in questo ragionamento sta nel ritenere che l'argomento diagonale sia un vero algoritmo per definire le cifre di r , mentre, in effetti, noi soltanto “immaginiamo” di numerare tutti i numeri calcolabili dell'intervallo $[0, 1)$, senza in pratica poterlo fare “effettivamente”. Questo problema è stato formalizzato e risolto da Alan Turing nel suo famosissimo lavoro [57] del 1936. La idea fondamentale di Turing è di distinguere gli algoritmi sintatticamente corretti che *terminano* la loro esecuzione, da quelli che *non terminano* la loro esecuzione perché contengono un *loop* che viene ripetuto senza fine.⁶

8.2.7 Teorema di Turing

Rivisitiamo allora l'esempio di un numero reale r non-calcolabile della Sezione 8.2.6. Ricordiamo che la definizione di r è basata sul procedimento diagonale di Cantor. Per comodità useremo ora per i numeri reali la rappresentazione decimale e considereremo solo i numeri dell'intervallo $[0, 1)$. Abbiamo detto che ogni algoritmo (programma) è una stringa finita di simboli in qualche linguaggio di programmazione e che il loro insieme è numerabile. Numeriamo allora con un indice $n \geq 1$ *tutti* i possibili programmi di un dato linguaggio di programmazione che ammettono come input un intero ≥ 1 , ordinandoli prima parzialmente per

⁶ Lo pseudocodice di un programma che *non termina* la sua esecuzione è, ad esempio,

```

 $x \leftarrow 0$ 
WHILE  $x \geq 0$ 
 $x \leftarrow x + 1$ 
```

dimensione e ordinandoli poi totalmente usando l'ordine alfabetico per quelli di ugual dimensione. Consideriamo il programma n -esimo $\mathcal{P}_n$ della lista: se $\mathcal{P}_n$ non è sintatticamente corretto, oppure se è corretto ma per l'input $n \geq 1$ non termina la sua esecuzione, oppure se per l'input n termina la sua esecuzione e genera come output un numero reale con la n -esima cifra decimale *diversa da 4*, definiamo $r_n = 4$; altrimenti (se $\mathcal{P}_i$ è sintatticamente corretto, e per l'input n termina la sua esecuzione e la n -esima cifra dell'output è 4), definiamo $r_n = 5$. Il numero r le cui cifre binarie sono $(r_n)_{n \geq 1}$ non è calcolabile, perché la sua cifra decimale n -esima è diversa da quella (se esiste) del numero calcolato dal programma n -esimo.

Appurato che r non è calcolabile, Turing analizza il presunto “calcolo” di r . Tale “calcolo” prevede che si possa stabilire, ad esempio, se il programma numero 14, che potrebbe aver già generato 13 cifre, genererà prima o poi la 14-esima cifra decimale. Dovremmo quindi, per poter “calcolare” r , possedere un criterio algoritmico che ci dica se il programma sintatticamente corretto numero n genererà sicuramente, prima o poi, la n -esima cifra s_n dell'output (e allora vedremo cosa è questa cifra e definiremo r_n di conseguenza), oppure se girerà per sempre (nel qual caso dovremo porre $r_n = 4$). Se possedessimo un tale criterio algoritmico, allora avremmo l'assurdo di un numero simultaneamente calcolabile e non-calcolabile. Quindi, conclude semplicemente Turing, *un tale criterio algoritmico non può esistere*. È questo, essenzialmente, il contenuto del *Teorema di Indecidibilità di Turing* [57]. Intuitivamente:

Teorema 8.2.6 *Non esiste alcun algoritmo in grado di decidere se un programma termina la sua esecuzione, oppure no.*

Il paradosso che si crea applicando il procedimento diagonale di Cantor all'insieme numerabile dei numeri calcolabili viene quindi risolto dal Teorema di Indecidibilità di Turing. Esistono quindi numeri reali che non possono essere calcolati e non saranno mai calcolati dall'uomo.

8.2.8 Numeri definibili e Paradosso di Richard

La nozione di numero calcolabile non va confusa con quella di *numero definibile*. La nozione di numero definibile si riferisce ad un particolare linguaggio $\mathcal{L}$ (naturale, come l'italiano, o formalizzato), ed è legata al *Paradosso di Richard*, proposto nel 1905 dal matematico francese Jules Richard [50].

Dato un linguaggio $\mathcal{L}$, un numero reale è detto *definibile* (o $\mathcal{L}$ -definibile) se esiste una frase scritta con un numero finito di caratteri del linguaggio $\mathcal{L}$ che lo definisce. Fissiamo, ad esempio, $\mathcal{L} = \text{italiano}$; il numero π è $\mathcal{L}$ -definibile, in quanto, ad esempio, è definito dalla frase

$\mathcal{P} = \text{“Sia } \pi \text{ il rapporto fra la lunghezza della circonferenza ed il diametro di un cerchio di raggio uguale a uno”}$

Siccome il numero di caratteri dell'italiano (esteso ai simboli matematici usuali) è finito, abbiamo che la totalità delle possibili definizioni è un insieme numerabile, e quindi i numeri definibili costituiscono un sottoinsieme numerabile $\mathcal{D}$ di $\mathbb{R}$. Possiamo allora elencare le rappresentazioni decimali dei numeri definibili, ordinandole secondo l'ordine alfabetico delle corrispondenti definizioni, ed applicare il procedimento diagonale di Cantor alle loro mantisse, ottenendo un numero reale r che *non è definibile* perché non è nella lista. Come ammette un filosofo laureato in matematica, Giulio Giorello [21]:

...: resta aperta infatti la questione di quale senso si possa attribuire alla esistenza di numeri reali non definiti da alcuna legge (ovvero la cui legge non può venire enunciata in un numero finito di parole).

D'altra parte, vi è chi ravvisa nella procedura diagonale che ha prodotto il numero r proprio la sua definizione: si avrebbe quindi un paradosso, in quanto r risulterebbe sia definibile che non-definibile (Paradosso di Richard). Possiamo proporre una possibile soluzione del Paradosso di Richard ripetendo il ragionamento che ha portato al Teorema di Indecidibilità di Turing. Scegliamo il linguaggio $\mathcal{L} = \text{italiano}$ (esteso ai simboli matematici usuali). Numeriamo con un indice $i \geq 1$ tutte le possibili frasi della lingua italiana (con un numero finito di caratteri, inclusi gli spazi ed i simboli matematici), ordinandole prima per lunghezza e poi ordinando quelle di ugual lunghezza per ordine alfabetico. Consideriamo la n -esima frase della lista: se non è sintatticamente corretta, o se non definisce un numero, o se definisce un numero con la n -esima cifra decimale della sua mantissa *diversa da 4*, definiamo $r_n = 4$; altrimenti (se la n -esima frase della lista è sintatticamente corretta e definisce un numero con la n -esima cifra decimale della sua mantissa *uguale a 4*), definiamo $r_n = 5$. Il numero $r = 0.r_1r_2r_3r_4\dots$ non è definibile, perché la sua cifra decimale n -esima è diversa da quella del numero definito dalla frase n -esima (se questa definisce un numero). Tuttavia abbiamo appena fornito in italiano la definizione di r . A questo assurdo giungiamo se ammettiamo che sia possibile decidere, per ogni frase della lingua italiana, se questa sia o non sia

la definizione di un numero. Se ammettiamo che esista un criterio con cui decidere se una frase definisce un numero o no, giungiamo all'assurdo di un numero simultaneamente definibile e non definibile. Quindi un tale criterio non esiste, ossia, *non può esistere in italiano (né in altre lingue) la definizione di definibilità*.

In effetti, se per ogni frase della lingua italiana fosse possibile decidere se è la definizione di un numero oppure no, basterebbe applicare tale criterio decisionale alla frase:

$\mathcal{G}$ = “Sia n il minimo numero intero pari maggiore di 2 che non è somma di due numeri primi”

per risolvere la Congettura di Goldbach:⁷ se $\mathcal{G}$ è una definizione, la Congettura di Goldbach è falsa in quanto il numero n definito da $\mathcal{G}$ è un controesempio, mentre se $\mathcal{G}$ non è una definizione, la Congettura di Goldbach è vera.

8.2.9 La “definizione” di successione

Il Paradosso di Richard induce una riflessione anche sulla definizione “discorsiva” di successione $a : \mathbb{N} \rightarrow A$ a valori in un insieme A (non vuoto). Infatti, la prassi didattica tradizionale definisce una successione a valori in A come una legge che ad ogni numero naturale n associa un elemento $a_n \in A$ (vedi ad es. [22], p. 52). Se per “legge” si deve intendere una formula in un linguaggio formalizzato, oppure un’espressione in un linguaggio naturale come l’italiano o l’inglese, ecc., allora il ragionamento di Richard mostra che, quale che sia l’insieme A , le successioni a valori in A sono al più in quantità numerabile. In particolare, sarebbe numerabile l’insieme delle successioni di numeri razionali, e quindi il modello di Méray–Cantor non sarebbe in grado di descrivere tutti i numeri reali, proprio a causa dell’argomento diagonale dello stesso Cantor (Sezione 8.2.3). Per risolvere questo problema Jules Richard ([50], vedi anche [21]) fu costretto ad ammettere la necessità di considerare successioni senza leggi, cioè “successioni la cui legge non può venir enunciata con un numero finito di parole”.

Questa difficoltà “linguistica” è tecnicamente superata non appena si definisca una funzione $f : A \rightarrow B$ come un sottoinsieme $f \subseteq A \times B$ tale che

⁷ La Congettura di Goldbach nasce nel 1742 in una lettera scritta dal matematico prussiano Christian Goldbach a Leonardo Eulero, e dice che *ogni numero pari maggiore di 2 può essere scritto come somma di due numeri primi* (non necessariamente distinti). È tuttora (2008) irrisolta.

per ogni $a \in A$ esiste uno ed un solo $b \in B$ (indicato con $b = f(a)$) tale che $(a, b) \in f$, in quanto gli assiomi generalmente riconosciuti della teoria degli insiemi garantiscono che tutti gli oggetti necessari a questa definizione *esistono*. In questo modo, l'insieme di tutte le funzioni $f : A \rightarrow B$ *esiste*, anche se, per il ragionamento di Richard, non tutte possono essere descritte o definite a parole.

8.3 Numeri reali e misura

8.3.1 Insiemi trascurabili

Ricordiamo che:

Definizione 8.3.1 *Un insieme $E \subseteq \mathbb{R}$ è trascurabile se per ogni $\epsilon > 0$ esiste una successione $([a_i, b_i])_{i \geq 0}$ di intervalli tale che si abbia $E \subseteq \bigcup_{i \geq 0} [a_i, b_i]$, con la somma della serie $\sum_{i=0}^{\infty} (b_i - a_i) \leq \epsilon$.*

Nell'ambito della teoria della misura di Lebesgue, gli insiemi trascurabili come qui definiti sono esattamente gli insiemi di misura (di Lebesgue) nulla. Vediamo allora che:

Proposizione 8.3.2 *Un insieme numerabile E di numeri reali è un insieme trascurabile.*

DIMOSTRAZIONE. Esiste una numerazione $(x_i)_{i \geq 0}$ di E . Dato $\epsilon > 0$, basta considerare per ogni $i \geq 0$ l'intervallo $[a_i, b_i]$ con $a_i = x_i - \epsilon/2^{i+2}$, $b_i = x_i + \epsilon/2^{i+2}$, di ampiezza $b_i - a_i = \epsilon/2^{i+1}$. Infatti, $\sum_{i=0}^{\infty} (b_i - a_i) = \sum_{i=0}^{\infty} \epsilon/2^{i+1} = \epsilon$. $\square$

Proposizione 8.3.3 *L'intervallo $[0, 1] \subseteq \mathbb{R}$ non è un insieme trascurabile.*

DIMOSTRAZIONE. Supponiamo per assurdo che lo sia. Dato $\epsilon = 0.5$, esiste una successione $([a_i, b_i])_{i \geq 0}$ di intervalli tale che si abbia $[0, 1] \subseteq \bigcup_{i \geq 0} [a_i, b_i]$, con la somma della serie $\sum_{i=0}^{\infty} (b_i - a_i) \leq 0.5$. Possiamo “allargare” del 10% gli intervalli, ottenendo una successione $((c_i, d_i))_{i \geq 0}$ di intervalli *aperti* tale che si abbia $[0, 1] \subseteq \bigcup_{i \geq 0} (c_i, d_i)$, con la somma della serie $\sum_{i=0}^{\infty} (d_i - c_i) \leq 0.55$ (basta porre $c_i = a_i - \delta/2$, $d_i = b_i + \delta/2$, con

$\delta = 0.1 \cdot (b_i - a_i)$). Per il Teorema di Heine–Borel 5.3.1, l'intervallo $[0, 1]$ è contenuto nella unione di un numero finito di intervalli, le cui lunghezze hanno somma al massimo pari a 0.55, evidentemente un assurdo. $\square$

Abbiamo quindi una seconda dimostrazione del fatto che un insieme numerabile (come ad esempio l'insieme $\mathbb{Q}$ dei numeri razionali, l'insieme $\mathcal{A}$ dei numeri algebrici, l'insieme $\mathcal{T}$ dei numeri calcolabili) non può esaurire tutto l'intervallo $[0, 1]$, e quindi esistono numeri irrazionali, esistono numeri trascendenti, esistono numeri non-calcolabili. Ma questa volta abbiamo ottenuto molto di più rispetto alla sola esistenza: essendo che i numeri razionali, gli algebrici, ed i calcolabili costituiscono insiemi trascurabili, possiamo mostrare che i numeri irrazionali, i trascendenti ed i numeri non-calcolabili costituiscono la “stragrande maggioranza” dei numeri reali. Per far capire questo, utilizzeremo la seguente:

Metafora del nastro adesivo Immaginiamo di avere di fronte a noi la retta reale $\mathbb{R}$. Consideriamo un sottoinsieme $E \subseteq \mathbb{R}$ numerabile. Sia $(x_i)_{i \geq 1}$ una numerazione di E . Immaginiamo di prendere un pezzo di nastro adesivo lungo (o corto) $\epsilon = 10$ cm, e di tagliarlo in due, ottenendo due spezzoni di 5 cm ciascuno, con uno dei quali “copriamo” il primo numero x_1 di E . Tagliamo in due il restante spezzone da 5 cm, ottenendo due spezzoni di 2.5 cm ciascuno, con uno dei quali “copriamo” il secondo numero x_2 . Tagliamo in due il restante spezzone da 2.5 cm, ottenendo due spezzoni di 1.25 cm ciascuno, con uno dei quali “copriamo” il terzo numero x_3 ; e così via... Possiamo così “coprire” con soli 10 cm di nastro adesivo *tutti* i numeri di E (oltre a vari numeri che non sono in E , “troppo vicini” a qualcuno degli x_i): tutti i numeri della retta reale che non sono stati “coperti” da qualche pezzo di nastro adesivo sono numeri del complementare $\mathbb{R} \setminus E$.

Ad esempio, non solo quindi esistono numeri reali non calcolabili, ma questi sono “molti di più” rispetto a quelli calcolabili. Utilizzando la metafora del nastro adesivo, bastano 10 cm di nastro per coprire sulla retta tutti i numeri reali calcolabili. Questo conferma la assoluta delicatezza, anche in termini filosofici, della nozione di numero reale. Per un approfondimento sui temi di questa Sezione, rimandiamo a [37].

8.3.2 L'insieme di Cantor

Avendo dimostrato che un insieme numerabile di numeri reali è un insieme trascurabile (Proposizione 8.3.2), ci possiamo chiedere se gli insiemi numerabili siano gli unici ad essere trascurabili: questo è falso. Presentiamo infatti in questa sezione l'*insieme (triadico) di Cantor* T_∞ , che è un insieme trascurabile pur avendo la stessa cardinalità dell'intervallo $[0, 1]$, ovvero la cardinalità del continuo. Consideriamo l'intervallo $T_0 = [0, 1] \subseteq \mathbb{R}$, e “rimuoviamo” da questo l'intervallo aperto $(\frac{1}{3}, \frac{2}{3})$ ottenendo

$$T_1 = [0, 1] \setminus (\frac{1}{3}, \frac{2}{3}) = [0, \frac{1}{3}] \cup [\frac{2}{3}, 1]$$

Rimuoviamo da T_1 i due “terzi di mezzo” $(\frac{1}{9}, \frac{2}{9})$ e $(\frac{7}{9}, \frac{8}{9})$ ottenendo

$$T_2 = [0, \frac{1}{9}] \cup [\frac{2}{9}, \frac{1}{3}] \cup [\frac{2}{3}, \frac{7}{9}] \cup [\frac{8}{9}, 1]$$

Induttivamente, T_n è definito rimuovendo da T_{n-1} i “terzi di mezzo”, e si definisce allora

$$T_\infty = \bigcap_{n \geq 0} T_n$$

Si vede subito che T_n è unione disgiunta di 2^n intervalli, ciascuno dei quali ha ampiezza $(\frac{1}{3})^n$; quindi la misura di T_n è $(\frac{2}{3})^n$, ed essendo che, per ogni $\epsilon > 0$, si ha $(\frac{2}{3})^n \leq \epsilon$ non appena sia $n \geq \log(\epsilon)/\log(\frac{2}{3})$, abbiamo che:

Proposizione 8.3.4 *L'insieme (triadico) di Cantor T_∞ è trascurabile.*

Tuttavia questo insieme *non* è numerabile, anzi, ha la stessa cardinalità dell'intervallo $[0, 1]$. Per provare ciò, osserviamo che

Proposizione 8.3.5 *L'insieme (triadico) di Cantor T_∞ è costituito da tutti e soli i numeri reali dell'intervallo $[0, 1]$ che ammettono una rappresentazione in base $b = 3$ senza la cifra ‘1’.*

DIMOSTRAZIONE. Ovviamente le cifre utilizzate nella rappresentazione in base $b = 3$ sono ‘0’, ‘1’ e ‘2’. Quindi, sono proprio i numeri dei “terzi di mezzo rimossi” ad avere la cifra ‘1’ ad un certo punto della loro rappresentazione in base $b = 3$. Ad esempio, i numeri dell'intervallo aperto $(\frac{1}{3}, \frac{2}{3})$, scritti in base $b = 3$, iniziano necessariamente con $0.1 \dots$, mentre i numeri dell'intervallo aperto $(\frac{7}{9}, \frac{8}{9})$, scritti in base $b = 3$, iniziano

necessariamente con $0.21\dots$ ⁸ Gli estremi invece appartengono a T_∞ , in quanto, ad esempio, $\frac{8}{9} = 0.22\overline{0}$, mentre $\frac{7}{9} = 0.21\overline{0}$ si può rappresentare anche senza la cifra ‘1’ come $\frac{7}{9} = 0.20\overline{2}$. $\square$

Dimostriamo ora che

Proposizione 8.3.6 *L'insieme (triadico) di Cantor T_∞ non è numerabile.*

DIMOSTRAZIONE. Le rappresentazioni in base $b = 3$ dei numeri reali dell'intervallo $[0, 1]$ senza la cifra ‘1’ costituiscono l'insieme $\{0, 2\}^\mathbb{N}$ di tutte le successioni dei due simboli zero ‘0’ e due ‘2’. Per l'argomento diagonale di Cantor, questo insieme non è numerabile. $\square$

8.3.3 Numeri normali

Un altro esempio, ma molto più complesso, di insieme trascurabile è dato dall'insieme dei numeri reali dell'intervallo $[0, 1]$ che *non* sono “normali”. Consideriamo un numero reale $x \in [0, 1]$, e sia $d_n(x)$ la n -esima cifra binaria di x , ossia

$$x = 0.d_1(x)d_2(x)d_3(x)d_4(x)\dots d_n(x)\dots_2 = \sum_{j=1}^{\infty} d_j(x)2^{-j}$$

Diremo che

Definizione 8.3.7 *Il numero reale x è normale (in base $b = 2$) se le frequenze asintotiche relative degli ‘0’ e degli ‘1’ nello sviluppo binario sono entrambe $= \frac{1}{2}$; formalmente, x è normale se*

$$\lim_{n \rightarrow \infty} \frac{d_1(x) + d_2(x) + \dots + d_n(x)}{n} = \frac{1}{2}$$

Ad esempio, $x = 2/3$ è normale in base 2, essendo $0.1010101010\dots_2 = 0.\overline{10}_2$ il suo sviluppo binario. Invece $x = 2/7$ non è normale in base 2, essendo $0.01001001001001001\dots_2 = 0.\overline{0100}_2$ il suo sviluppo binario. La definizione di normalità può essere estesa ad una base generica $b > 2$ in modo ovvio. Si osservi però che un numero x che è normale in base 2

⁸ Infatti, in base $b = 3$ si ha $0.1\dots = \frac{1}{3} + \dots$, mentre $0.21\dots = \frac{2}{3} + \frac{1}{9} + \dots = \frac{7}{9} + \dots$

può non essere normale in altre basi: ad esempio, $x = 2/3$ è normale in base 2, ma non lo è in base 10, essendo $0.6666666 \dots = 0.\overline{6}$ il suo sviluppo decimale.

Émile Borel dimostrò nel 1909 [8] che l'insieme dei numeri dell'intervallo $[0, 1]$ che *non* sono normali è un insieme trascurabile (Teorema di Borel, cf. [44]). Questo Teorema di Borel è un caso particolare della Legge Forte dei Grandi Numeri del calcolo delle probabilità.

8.4 Numeri reali e scienze cognitive

di Michela Maschietto

*Die ganzen Zahlen hat der liebe Gott gemacht;
alles andere ist Menschenwerk.^a*

LEOPOLD KRONECKER

^a Il buon Dio ha creato i numeri naturali; tutto il resto è opera dell'uomo.

I numeri naturali, cioè gli elementi di $\mathbb{N}$, sono da molto tempo oggetto di studio e interesse non solo dei matematici, come ben testimonia la storia della matematica, ma anche ad esempio dei ricercatori di neuroscienze, di psicologia, di neuropsicologia, $\dots$. Numerose sono le ricerche che si occupano della percezione della numerosità, del processo di conteggio e della rappresentazione dei numeri. Come riferimento bibliografico, possiamo citare le ricerche condotte da Stanislas Dehaene [17] e da Brian Butterworth [10]. In generale, tali ricerche hanno evidenziato meccanismi cerebrali di valutazione della numerosità, presenti non solo negli esseri umani ma anche in alcune scimmie antropomorfe. Sebbene con approcci diversi, la loro ricerca è volta a mostrare che il cervello degli esseri umani è dotato di circuiti atti a individuare e riconoscere le numerosità nel mondo circostante e che su questi circuiti si fonda la conoscenza matematica. Fra questi vi è ad esempio il “subitizing”, corrispondente alla capacità di percepire la numerosità di un insieme visivo di oggetti (quattro o cinque) immediatamente, senza mettere in atto processi di conteggio. Oltre

cinque oggetti, occorre attuare una qualche strategia di conteggio, soprattutto se la posizione spaziale degli oggetti non si presenta ordinata. Lo studio di soggetti aventi subito danni cerebrali (condotto nell'ambito della neuropsicologia cognitiva) ha permesso di studiare il legame tra capacità numeriche e altre capacità, come quelle linguistiche e di rappresentazione. Inoltre ha consentito di mappare alcune zone della corteccia cerebrale, attivate oppure no a seconda del compito assegnato.

I risultati di queste ricerche possono essere utili alla didattica della matematica, in quanto portano l'attenzione a ciò che si può considerare "innato", cioè biologicamente presente negli esseri umani, in particolare nei bambini, rispetto a ciò che può essere sollecitato o frutto di un processo di apprendimento/insegnamento. Ad esempio, per andare oltre cinque oggetti si ha bisogno di una serie di strumenti concettuali, che sono stati sviluppati dagli uomini nella storia. Ad esempio, si possono considerare le dita, le parole che esprimono i numeri e i simboli che li rappresentano. Questi strumenti possono diventare accessibili grazie a processi ben precisi, come quelli di apprendimento e insegnamento. Essi permettono anche di andare oltre la numerosità e cogliere le costruzioni di numeri che da questa si allontanano (per esempio già i numeri interi, cioè gli elementi di $\mathbb{Z}$). Come appare dai seppur pochi elementi forniti, nelle ricerche sopra citate non si considerano i numeri reali. Oltre a ragioni storiche e culturali, dal punto di vista dello sviluppo concettuale, i numeri reali hanno una struttura assolutamente diversa da quella della numerosità, non avendo un analogo diretto nel cervello (cioè non vi è un'immagine intuitiva come per i numeri naturali).

8.4.1 La scienza cognitiva della matematica

In tempi recenti, alcuni ricercatori hanno iniziato ad investigare la matematica con gli strumenti delle scienze cognitive da una diversa prospettiva. Tra questi, il lavoro che ha suscitato maggiore interesse è stato quello di George Lakoff e Rafael E. Núñez [35] *Where Mathematics Comes From*, principalmente per due ragioni. La prima è il tipo di analisi che viene condotta e l'ambizione del libro stesso, che si propone di "fondare" una scienza cognitiva della matematica, in cui componenti fondamentali

sono l’“embodiment”⁹ ed il pensiero metaforico.¹⁰ La seconda è legata alle reazioni di alcuni matematici sui contenuti matematici presenti nel testo e sulla loro presentazione (ad esempio, alcune “sviste”, rilevate ad esempio nella recensione scritta da Bonnie Gold per la Mathematical Association of America¹¹). Nel tentativo di andare oltre le prime critiche al testo, i ricercatori in educazione matematica Martin Schiralli e Nathalie Sinclair [53] hanno cercato da una parte di mettere in evidenza alcuni punti deboli di un tale approccio, dall’altra di proporre un loro approccio alla questione da dove viene la matematica, tenendo conto delle ricerche condotte in didattica della matematica. Un primo punto debole rilevato è la mancanza di un’adeguata esplicitazione dell’accezione con cui il termine “matematica” è usato nel libro: in esso non viene differenziato il termine “matematica” in relazione ad attività come apprendere, fare o usare la matematica. Un secondo punto debole della ricerca riguarda la metodologia seguita: affrontare la questione da dove viene la matematica richiede non solo l’analisi delle idee matematiche (obiettivo dichiarato del libro), ma anche un’indagine empirica delle differenze individuali e culturali presenti nell’apprendere, fare o usare la matematica. Al libro sono poi seguiti altri lavori, in cui si precisano sue parti [45] e si cerca di affrontare la questione del legame con la didattica della matematica [46].

La ricerca di Lakoff e Núñez si interessa allo studio dell’origine delle idee matematiche: come sono concettualizzate e che cosa significano. È condotta all’interno della problematica delle scienze cognitive riguardanti lo studio della struttura concettuale del pensiero umano. Un assunto di partenza di tale studio è che gli oggetti della matematica sono astrazioni mentali e come tali non possono essere percepiti attraverso i sensi. Tuttavia, la loro analisi è volta a mostrare che la maggior parte degli oggetti tecnici astratti idealizzati in matematica sono creati mediante

⁹ Secondo cui le caratteristiche del nostro corpo e del nostro cervello condizionano la struttura dei concetti che sono costruiti.

¹⁰ Secondo cui i concetti astratti sono, nella maggior parte dei casi, concettualizzati in termini concreti. Il meccanismo cognitivo soggiacente è chiamato *metafora concettuale*. Una metafora concettuale è una mappatura che conserva la struttura delle inferenze di un dominio sorgente quando è proiettata in un dominio bersaglio. In tal modo, il dominio bersaglio è compreso in termini di relazioni tra elementi nel dominio sorgente. Le metafore sono espresse nella forma “ A è B ”: la mappatura è $B \longrightarrow A$, con B dominio sorgente ed A dominio bersaglio.

¹¹ Cf. www.maa.org/reviews/wheremath.html.

meccanismi cognitivi che estendono la struttura dell'esperienza corporea preservando l'organizzazione delle inferenze dei campi d'esperienza [45]. In altri termini, la concettualizzazione in matematica si appoggia su meccanismi cognitivi che sono caratteristici dell'attività umana quotidiana e non specifici della sola matematica. L'attenzione è portata sulla messa in evidenza delle capacità “embodied” che permettono di passare dalle abilità numeriche innate alla comprensione della matematica. Gli autori distinguono tra matematica pura e idee matematiche e sono interessati a studiarne le relazioni.

Gli strumenti usati per l'analisi delle idee matematiche da una parte sono presi dalla linguistica cognitiva (in particolare, dalla semantica cognitiva); dall'altra derivano dalle ricerche che si occupano dell'analisi dei gesti in relazione allo sviluppo del pensiero. Il primo tipo di strumento ha permesso di individuare metafore e metonimie¹² nei discorsi sulla matematica, mentre il secondo ha sottolineato come le costruzioni linguistiche repertoriate, ad esempio, nei testi matematici non siano espressioni “morte”, ma caratterizzano il pensiero e la comunicazione delle idee matematiche.

Consideriamo un esempio del tipo di indagine e analisi condotti in [45]. Nel testo *What is mathematics?* di Richard Courant ed Herbert Robbins [15] è stata considerata una parte riguardante il limite, ovvero la somma, di serie numeriche. Ecco l'argomentazione proposta:

Nel caratterizzare il limite di serie infinite, Courant e Robbins scrivono:

“il comportamento di s_n si descrive dicendo che la somma s_n si avvicina al limite 1 quando n tende all'infinito, e scrivendo

$$1 = \frac{1}{2} + \frac{1}{2^2} + \frac{1}{2^3} + \frac{1}{2^4} + \dots" \quad (\text{p. 64}^{13}, \text{ nostra enfasi})$$

Strettamente parlando, questo enunciato si riferisce ad una successione s_n di somme parziali (numeri reali) discrete e prive di moto, corrispondenti a valori discreti, crescenti e privi di moto assunti da n nell'espressione $\frac{1}{2^n}$, con n un numero naturale. Se si esamina questo enunciato, si può vedere che esso descrive alcuni fatti sui numeri e sul risultato di operazioni sui numeri, ma che

¹² Una metonimia concettuale è un meccanismo cognitivo che permette di concepire una parte di un intero per l'intero stesso [45].

non c'è il benché minimo moto coinvolto. Nessuna entità si sta veramente avvicinando o tendendo a qualche cosa. Allora, perché Courant e Robbins (o i matematici in generale) usano un linguaggio dinamico per esprimere proprietà statiche di entità statiche? E che cosa significa dire che “la somma s_n si avvicina”, quando di fatto una somma è semplicemente un numero fissato, un risultato di un'operazione di addizione?¹⁴

L'enunciato evidenzia la presenza di riferimenti a un qualche moto. L'analisi proposta è ricondotta a quella della struttura dei numeri reali. Nel sistema di assiomi di $\mathbb{R}$ non vi è riferimento ad una somma “che si avvicina” ad un numero o ad un numero “che tende” all'infinito. Sono gli strumenti presi dalle scienze cognitive a fornire una chiave interpretativa dell'enunciato soprariportato. In esso (come in altri che si possono trovare nei testi matematici), il moto emerge in modo metaforico dai successivi valori presi da n nella successione. L'analisi linguistica mette in evidenza una serie di metafore e metonimie, che non possono essere considerate come se si trattassero di elementi presenti solo in uno stadio di preformalismo. Esse sono di fatto costitutive delle idee “embodied” che rendono le idee matematiche possibili. Il riferimento al moto presente nel testo di Courant e Robbins è fatto risalire a un meccanismo chiamato “fictive motion” (studiato da L. Talmy [55, 56]) e così definito: il moto fittizio è un fondamentale meccanismo cognitivo “embodied” attraverso il quale inconsciamente (e senza sforzo) si concettualizzano entità statiche in termini dinamici. Ad esempio, questo meccanismo si riscontra nell'espressione “l'equatore *attraversa* molti paesi”. Quando si parla di moto fittizio, vi è sempre un tracciatore (l'agente che si muove) e un paesaggio (lo spazio in cui il tracciatore si muove). Il tracciatore può essere un oggetto reale o un'entità immaginaria. In matematica, tuttavia, il tracciatore ha sempre una componente metaforica.

L'analisi di altri esempi porta a concludere che la struttura delle idee matematiche, e la sua organizzazione di inferenze, è più ricca e più dettagliata della struttura di inferenze fornita dalle definizioni formali e dai metodi assiomatici. Le definizioni formali e gli assiomi né formalizzano interamente né generalizzano i concetti umani.

La distinzione tra idee matematiche e matematica pura risulta ancor

¹⁴ Nostra traduzione.

più evidente nell'analisi della continuità di funzione, in cui gli autori distinguono tra continuità naturale e continuità di Cauchy–Weierstrass.¹⁵ La continuità naturale corrisponde all'idea di continuità presente nei discorsi quotidiani. Essa si basa sul fatto che un processo è considerato continuo se si sviluppa senza salti, interruzioni o cambiamenti repentini. Secondo gli autori, è questa l'idea presente negli scritti dei matematici del XVII secolo (per esempio in Euler). In accordo con tale continuità, l'idea di funzione continua presenta le seguenti caratteristiche:

- una funzione continua è determinata da un movimento continuo;
- il movimento implica una direzionalità per la funzione;
- il movimento ha come risultato una linea che non presenta salti;
- il movimento implica che vi è qualcosa che si sposta.

Dal punto di vista cognitivo, si ritrova il meccanismo del moto fittizio. La definizione di continuità è presente per Lakoff e Núñez come risposta ad un'esigenza di rigore nel processo di evoluzione della matematica nel XIX secolo. Lo studio di questa continuità s'inserisce nello studio del processo d'aritmetizzazione della matematica, che gli autori considerano all'interno di un programma di discretizzazione dello spazio. Questo corrisponde a un processo di concettualizzazione dello spazio naturalmente continuo e della continuità naturale in termini discreti, ad esempio in termini di punti e numeri.¹⁶ Nell'analisi, gli autori portano l'attenzione sul fatto che la definizione basata sui punti è una costruzione matematica tecnica, che implica un formidabile lavoro di costruzione di metafore. In particolare, la definizione della continuità di funzione che ne risulta si basa su tre metafore concettuali:

1. una linea è un insieme di punti;
2. la continuità naturale è assenza di buchi;
3. tendere verso un limite è preservare la prossimità rispetto a un punto.

¹⁵ La continuità secondo Cauchy–Weierstrass di una funzione $f : \mathbb{R} \rightarrow \mathbb{R}$ in un punto è quella legata alla definizione di limite data con i classici ϵ e δ .

¹⁶ In un tale processo, un ruolo centrale è ricoperto dalla metafora “uno spazio è un insieme di punti”.

Nei libri di testo di matematica (ad esempio, in quelli per la scuola secondaria), i due tipi di continuità sono in generale considerati come strettamente legati: spesso il primo è usato come argomento intuitivo per introdurre il secondo. Tuttavia, il confronto tra i due tipi di continuità ne mette in evidenza le differenze, in quanto esse sono associate a due strutture concettuali differenti. La continuità naturale sembra essere una costruzione “primitiva” per la presenza del riferimento al movimento, mentre la continuità di Cauchy–Weierstrass è una costruzione molto più complessa. Secondo gli autori, le conseguenze didattiche di quest’analisi sono importanti: non si può considerare il primo tipo di continuità come intuitiva e introdurre in seguito la definizione di Cauchy–Weierstrass in termini di sua formalizzazione. L’insegnamento dovrebbe prendere in considerazione il fatto che gli allievi si trovano effettivamente confrontati a due costruzioni concettuali diverse.¹⁷

Come è stato precedentemente detto, l’analisi delle idee matematiche si fonda sul ricorso a strumenti sia linguistici che gestuali. L’analisi dei gesti, poco presente nel libro *Where Mathematics Comes From*, ha acquistato maggiore importanza nei lavori successivi [45, 46], andando a prendere in considerazione uno dei punti deboli evidenziati da Schiralli e Sinclair [53]. L’osservazione di lezioni di matematica sui limiti, successioni, continuità ha evidenziato che quelle espressioni linguistiche sono accompagnate da gesti che enfatizzano gli aspetti dinamici. L’analisi dei gesti conferma quanto trovato nell’analisi linguistica, cioè che *il linguaggio formale in matematica non può catturare la complessità dell’organizzazione delle inferenze delle idee matematiche. Il lavoro della scienza cognitiva “embodied” è quello di caratterizzare la ricchezza delle idee matematiche.*

8.4.2 Invarianti e fondamenti della matematica

Lakoff e Núñez non sono tuttavia gli unici studiosi che hanno iniziato ad occuparsi del legame tra matematica (non solo quella parte riguardante i numeri naturali) e cognizione. Nello stesso periodo, da un altro punto di vista, il gruppo francese “Géométrie et Cognition”¹⁸ si è interessato

¹⁷ Questa rottura è discussa anche nei lavori di ricerca in didattica della matematica.

¹⁸ Gruppo con J. Petitot e B. Teissier in seno all’École Normale Supérieure di Parigi. Sito web: www.di.ens.fr/~longo/geocogni.html.

allo studio del legame tra geometria e cognizione, partendo dalla questione dei fondamenti della matematica. Anche se il *focus* del gruppo è la geometria e le relazioni spaziali, le riflessioni condotte contengono interessanti elementi di analisi. Lo scopo di tale lavoro è duplice: si cerca da una parte di integrare lo studio dei fondamenti della matematica e lo studio delle funzioni cognitive, dall'altra di elaborare dei modelli matematici di tali funzioni. L'articolo "The mathematical continuum: from intuition to logic" [38] presenta interessanti elementi di analisi dei legami tra intuizione e formalizzazione del continuo. Tale studio è anche motivato dal fatto che la costruzione del concetto di continuo ha coinvolto diversi matematici ed ha incontrato non poche difficoltà. L'autore si propone di affrontare la questione del continuo dalla prospettiva di cercare di ottenere un fondamento per la conoscenza matematica come parte del nostro modo di interpretare e ricostruire il mondo, non solo in termini di una pura investigazione logica della matematica.

Basandosi sui lavori di Weyl [60] e di Husserl, l'analisi evidenzia in tre livelli: quello dell'intuizione, quello della matematica (costruzione di Cantor–Dedekind) e quello della logica. La caratteristica importante di tale approccio è l'importanza delle influenze reciproche tra i vari livelli.

Nel primo livello, l'autore considera le diverse esperienze legate al continuo: la percezione del tempo, il movimento di un corpo, la traccia di una matita su un foglio. La posizione dell'autore è che l'intuizione del continuo è costruita partendo da elementi comuni o stabili, da invarianti che emergono da una pluralità di atti di esperienza. Nel livello matematico, l'autore esamina il passaggio tra le esperienze del primo livello alle costruzioni del continuo, integrando una componente epistemologica e una cognitiva. L'analisi epistemologica evidenzia le scelte possibili ad un certo momento storico (tra l'approccio di Leibniz e quello in termini di limite) e le ragioni della scelta effettuata. L'analisi cognitiva cerca di fornire una spiegazione della costruzione della definizione in termini di invarianti soggiacenti e della loro presa in considerazione nella concettualizzazione del continuo. In tal modo, la definizione di Cantor–Dedekind risulta basata su tre invarianti:

- invarianza di scala: tutte le parti, anche le più piccole, di tempo, di una linea... conservano le stesse proprietà di una parte più grande.

Questo è legato al fatto che, secondo l'autore, i cambiamenti di scala non modificano le intuizioni del continuo;

- assenza di salti;
- assenza di buchi o di punti singoli.

Su questa costruzione, due importanti osservazioni sono proposte. La prima è che il continuo derivante direttamente delle esperienze è in qualche modo distrutto per raggiungere i punti ed in seguito ricostruito a partire dai punti. La seconda è che una tale costruzione ha fornito le regole di costruzione dei numeri reali. L'esempio dell'analisi del continuo permette all'autore di mostrare come la costruzione matematica sia basata sulle esperienze vissute e sull'individuazione di invarianti. Per queste ragioni essa non è arbitraria, anche se altre costruzioni sono possibili, come quella di Leibniz.

Tutte queste ricerche, condotte principalmente al di fuori del campo della ricerca in didattica della matematica, sono state discusse e approfondite in tale ambito da diversi ricercatori, in particolare nell'Università di Torino (cf. ad esempio [40]). I lavori in didattica della matematica hanno nel tempo integrato gli elementi di tali studi in un più ampio approccio semiotico ai processi di apprendimento/insegnamento della matematica. In tale approccio, è stato dato particolare rilievo allo studio dei gesti [1] (alcuni elementi sono presenti anche in [20] e [52]).

Appendice A

Postfazione

Le conversazioni avute con alcuni docenti che hanno avuto la possibilità di leggere in anteprima questo libro, ci hanno fatto capire che una delle domande che, più di altre, i lettori si porranno è quale sia il percorso da seguire per introdurre i numeri reali nei corsi di matematica di base in modo rigoroso e didatticamente efficace, sia che il docente adotti un approccio assiomatico sia che scelga un approccio costruttivo.

A nostro parere, i risultati e le discussioni dei Capitoli 2 e 5 mostrano che per l'approccio assiomatico la definizione di $\mathbb{R}$ preferibile è quella di un campo ordinato nel quale ogni sottoinsieme non-vuoto e superiormente limitato ha l'estremo superiore. Questo naturalmente richiede la definizione preliminare di estremo superiore, e qualche esercizio in proposito, ma si tratta di argomenti che, di norma, sono comunque svolti. Il vantaggio è che non serve dimostrare il teorema di esistenza dell'estremo superiore, dato che risulta vero per definizione. Il teorema di esistenza dell'estremo superiore è peraltro il primo e principale (e didatticamente difficile) obiettivo degli approcci assiomatici basati sulle proprietà di separazione; queste ultime vengono quasi sempre usate solo per provare questo teorema, per cui tanto vale farne a meno.

L'approccio costruttivo è generalmente evitato dai docenti, perché considerato troppo complicato. Questo è certamente vero per i modelli di Dedekind e di Méray–Cantor, presentati nei Capitoli 3 e 4. Tuttavia, proprio lo studio dettagliato di questi modelli ci indica che è possibile semplificare la costruzione di Dedekind nel modo accennato nella Sezione 2.6.4 (semirette di Russell) in modo da ottenere un modello più

“didattico”. In questo caso però, converrà introdurre i numeri negativi per ultimi, seguendo la strada $\mathbb{N} \rightarrow \mathbb{Q}^+ \rightarrow \mathbb{R}^+ \rightarrow \mathbb{R}$ invece della $\mathbb{N} \rightarrow \mathbb{Z} \rightarrow \mathbb{Q} \rightarrow \mathbb{R}$ che è algebricamente più elegante, ma geometricamente meno intuitiva. Accenniamo brevemente a questa possibilità. In luogo delle semirette considereremo dei *segmenti* e ciò è più vicino sia all’impostazione originale di Russell¹ che all’intuizione geometrica della misura delle grandezze lineari.

A.1 Un modello “didattico” di $\mathbb{R}$

A.1.1 Segmenti di $\mathbb{Q}^+$

Indichiamo con $\mathbb{Q}^+$ l’insieme dei numeri razionali positivi o nulli.

Definizione A.1.1 *Un segmento di $\mathbb{Q}^+$ è un sottoinsieme $A \subseteq \mathbb{Q}^+$ tale che:*

- (S1) $0 \in A$,
- (S2) A è superiormente limitato,
- (S3) per ogni $a' \in \mathbb{Q}^+$, da $a' < a \in A$ segue $a' \in A$.

Chiamiamo numeri reali (positivi o nulli) i segmenti di $\mathbb{Q}^+$, ed indichiamo il loro insieme con $\mathbb{R}^+$.

Si osservi che $\{0\}$ è un segmento di $\mathbb{Q}$. Ora i segmenti A di $\mathbb{Q}^+$ possono essere di due tipi: quelli per cui esiste $r = \sup A \in \mathbb{Q}$, e quelli privi di estremo superiore in $\mathbb{Q}$. I primi possono essere detti *segmenti di prima specie*, o *segmenti commensurabili*, e sono intervalli del tipo $[0, r)$ oppure del tipo $[0, r]$. I secondi possono essere detti *segmenti di seconda specie*, o *segmenti incommensurabili*. È evidente che se A è un segmento commensurabile, il numero razionale $r = \sup A$ può essere assunto come misura della “lunghezza” di A . Un esempio di segmento incommensurabile è quello che rappresenta la diagonale del quadrato di lato unitario, ossia:

$$D = \{a \in \mathbb{Q}^+ \mid a^2 < 2\} \tag{A.1}$$

¹ Infatti il termine originale usato da Bertrand Russell non è *semiretta*, ma *segmento*; scrive Russell “We shall define the real numbers as segments of the series of rational numbers, in order to be sure of their existence.” (cf. [58], p. 220).

A.1.2 Operazioni ed ordine fra segmenti di $\mathbb{Q}^+$

Se A, B sono segmenti di $\mathbb{Q}^+$ definiamo, in modo molto naturale:

$$A + B = \{a + b \mid a \in A, b \in B\} \quad (\text{A.2})$$

$$A \cdot B = \{a \cdot b \mid a \in A, b \in B\} \quad (\text{A.3})$$

$$A \leq B \iff A \subseteq B \quad (\text{A.4})$$

Proviamo ora che:

Proposizione A.1.2 *Se A, B sono segmenti commensurabili di $\mathbb{Q}^+$, anche gli insiemi $A + B$ e $A \cdot B$ sono segmenti commensurabili di $\mathbb{Q}^+$, e si ha:*

$$\sup(A + B) = \sup A + \sup B \quad (\text{A.5})$$

$$\sup(A \cdot B) = \sup A \cdot \sup B \quad (\text{A.6})$$

$$A \leq B \implies \sup A \leq \sup B \quad (\text{A.7})$$

DIMOSTRAZIONE. Proviamo per prima cosa che $A + B$ e $A \cdot B$ sono segmenti di $\mathbb{Q}^+$. Le (S1) ed (S2) sono evidenti ($0 = 0 + 0 = 0 \cdot 0$; inoltre, se a è una maggiorazione di A e b è una maggiorazione di B , $a + b$ è una maggiorazione di $A + B$ e $a \cdot b$ è una maggiorazione di $A \cdot B$). Proviamo la (S3) per $A + B$. Se $a' \in \mathbb{Q}^+$, ed è $a' < a + b$, con $a \in A, b \in B$, posto $c = a \cdot a' / (a + b)$, $d = b \cdot a' / (a + b)$, dato che $a' / (a + b) < 1$ si ha $c < a$, $d < b$ e per la (S3) applicata ad A e B si ha che $c \in A$ e $d \in B$, e quindi $c + d = a' \in A + B$. Per provare la (S3) per $A \cdot B$, conviene usare la rappresentazione

$$A \cdot B = \bigcup_{b \in B} \{a \cdot b \mid a \in A\}$$

È infatti immediato provare che, per $b > 0$ fissato, l'insieme $\{a \cdot b \mid a \in A\}$ è un segmento, e che l'unione arbitraria di segmenti è un segmento.

Proviamo ora la (A.5). Poniamo $a = \sup A$, $b = \sup B$. Per ogni $\varepsilon \in \mathbb{Q}$, $\varepsilon > 0$, esistono $x^* \in A$ e $y^* \in B$ tali che $x^* > a - \frac{\varepsilon}{2}$ e $y^* > b - \frac{\varepsilon}{2}$. Pertanto, $x^* + y^* > a + b - \varepsilon$. Poiché $z^* = x^* + y^* \in A + B$, è provato che $a + b = \sup(A + B)$.

Proviamo la (A.6). Poniamo $a = \sup A$, $b = \sup B$. Sia dato un arbitrario $\varepsilon' \in \mathbb{Q}$, $\varepsilon' > 0$. Scegliamo $\varepsilon \in \mathbb{Q}$ tale che

$$0 < \varepsilon < \min\{a + b, \varepsilon' / (a + b)\}$$

Esistono $x^* \in A$ e $y^* \in B$ tali che $x^* > a - \varepsilon$ e $y^* > b - \varepsilon$. Pertanto

$$x^* \cdot y^* > ab - \varepsilon \cdot (a + b - \varepsilon)$$

Ora da $0 < \varepsilon < a + b$ segue che $0 < a + b - \varepsilon < a + b$; quindi $1/(a + b) < 1/(a + b - \varepsilon)$ e allora $\varepsilon < \varepsilon'/(a + b) < \varepsilon'/(a + b - \varepsilon)$, e infine

$$x^* \cdot y^* > ab - \varepsilon'$$

Poiché $z^* = x^* \cdot y^* \in A \cdot B$, è provato che $a \cdot b = \sup(A \cdot B)$.

La (A.7) è immediata. □

L'applicazione

$$\sup : \{\text{segmenti commensurabili di } \mathbb{Q}^+\} \rightarrow \mathbb{Q}^+$$

che ad ogni segmento commensurabile di $\mathbb{Q}^+$, con estremo inferiore 0, fa corrispondere il suo estremo superiore, ossia la sua “lunghezza”, conserva quindi le operazioni di addizione e moltiplicazione e la relazione d'ordine. È quindi possibile identificare i segmenti commensurabili (numeri reali, o segmenti, di prima specie) con il loro estremo superiore, ossia con la misura (razionale) della loro “lunghezza”. Per esemplificare il caso dei segmenti incommensurabili, consideriamo la diagonale del quadrato di lato unitario:

$$D = \{a \in \mathbb{Q}^+ \mid a^2 < 2\} \tag{A.8}$$

È ragionevole pensare che alla fine la “lunghezza” di questa diagonale debba essere il numero reale $\sqrt{2}$. Ma evidentemente $\sqrt{2}$ “è” il segmento (incommensurabile) D . Quindi la “lunghezza” di un segmento incommensurabile S è il segmento stesso pensato come numero reale. In termini geometrici, la lunghezza di un segmento incommensurabile S è l'*insieme* di tutti i segmenti commensurabili “più corti” di S .

A.1.3 Proprietà delle operazioni e dell'ordine

Le verifiche delle proprietà algebriche delle operazioni e della relazione d'ordine si effettuano come nel modello di Dedekind. In una presentazione didattica, contenuta nel tempo, questi dettagli dimostrativi possono essere omessi.

A.1.4 Completezza

Proviamo ora il teorema di esistenza dell'estremo superiore per $\mathbb{R}^+$. Sia E un insieme non-vuoto e superiormente limitato di numeri reali positivi o nulli; ogni $x \in E$ è quindi un segmento costituito da numeri razionali positivi. Poniamo

$$\xi = \bigcup_{x \in E} x$$

ovvero formiamo l'unione di tutti i segmenti $x \in E$. Tale unione ξ è ancora un segmento di $\mathbb{Q}^+$, e dato che $x \subseteq \xi$ (ovvero $x \leq \xi$) per ogni $x \in E$, si ha che ξ è una maggiorazione di E in $\mathbb{R}^+$. Sia ora $\eta \in \mathbb{R}^+$, $\eta < \xi$; ciò significa che η è un segmento strettamente contenuto in $\bigcup_{x \in E} x$, e quindi esiste un razionale $r \in \bigcup_{x \in E} x$, $r \notin \eta$. Tale r deve appartenere ad un opportuno $x^* \in E$ e pertanto η è un segmento strettamente contenuto in x^* . Abbiamo allora verificato che in $\mathbb{R}^+$ si ha $\xi = \sup E$.

A.1.5 Conclusione della costruzione

Per concludere la costruzione dobbiamo introdurre i reali negativi: questo si fa come² nella classica estensione $\mathbb{N} \rightarrow \mathbb{Z}$ e si prolunga la moltiplicazione usando la regola dei segni. Abbiamo così ottenuto $\mathbb{R}$.

² In termini algebrici, si usa la simmetrizzazione di un monoide abeliano.

Bibliografia

- [1] F. Arzarello, L. Edwards, *Gesture and the Construction of Mathematical Meaning*, URL: www.emis.de/proceedings/PME29/PME29ResearchForums/PME29RFArzarelloEdwards.pdf
- [2] C. Arzelà, *Trattato di algebra elementare: ad uso dei licei*, Firenze: Succ. Le Monnier, 1880.
- [3] B. Banaschewski, On Proving the Existence of Complete Ordered Fields, *Amer. Math. Monthly*, 105(6), 1998, 548–551.
- [4] G. C. Barozzi, E. Gonzalez, *Calculus*, Edizioni Libreria Progetto, Padova 1969, Vol. 1.
- [5] L. Barthel, A computerized interactive approach to real numbers and decimal expansion, ICME 10 Proceedings, Copenhagen, 2004. URL: www.icme-organisers.dk/tsg12/papers/barthel-tsg12.pdf.
- [6] R. Beigel, Irrationality Without Number Theory, *Amer. Math. Monthly*, 98(4), 1991, 332–335.
- [7] V. Benci, *Lezioni di Analisi Matematica, esposte col metodo degli infinitesimi*, Pisa, Tipografia Editrice Pisana, 1998.
- [8] È. Borel, *Leçons sur la Théorie des Fonctions*, Gabay, Paris, 1950.
- [9] È. Borel, Les probabilités dénombrables et leurs applications arithmétiques, *Rend. Circ. Mat. Palermo* 27 (1909), 247–271.
- [10] B. Butterworth, *The mathematical brain*, Macmillan, 1999; trad. it.: C. Capararo, S. Mancini, e B. Isabella, *Intelligenza matematica*, Rizzoli, Milano, 1999.

- [11] G. Cantor, Über die Ausdehnung eines Satzes aus der Theorie der trigonometrischen Reihen, *Math. Annalen* 5(1872), 123–132.
- [12] G. J. Chaitin, *Teoria algoritmica della complessità*, Torino, Giappichelli, 2006.
- [13] A. Capelli, *Lezioni di algebra complementare: ad uso degli aspiranti alla licenza universitaria in scienze fisiche e matematiche*, Napoli, B. Pellerano, 1895.
- [14] P. M. Cohn, *Algebra*, Vol. 1, Aberdeen, John Wiley & Sons, 1974.
- [15] R. Courant, H. Robbins, I. Stewart, *What is Mathematics?: An Elementary Approach to Ideas and Methods*, 2nd ed., Oxford University Press US, 1996; trad. it.: L. Ragusa Gilli, *Che cos'è la matematica?*, 2^a ed., Bollati Boringhieri, Torino, 2005.
- [16] R. Dedekind, *Essenza e significato dei numeri. Continuità e numeri irrazionali*, Trad. dal tedesco e note storico-critiche di O. Zariski, Stock, Roma, 1926.
- [17] S. Dehaene, *La bosse des maths*, Éditions Odile Jacon, Paris, 1997; trad. it.: *Il pallino della matematica*, Mondadori, Milano, 2000.
- [18] D. Dikranjan, M. S. Lucido, *Aritmetica e algebra*, Napoli, Liguori, 2007.
- [19] M. Dolcher, *Elementi di analisi matematica*, Vol. I, Lint, Trieste, 1991.
- [20] F. Ferrara, *Acting and interacting with tools to understand Calculus Concepts*, Tesi di dottorato, Università di Torino, 2006.
- [21] G. Giorcello, *Transfinito e continuo*, tesi di laurea in matematica, 1971. URL:
www.tesionline.it/tesi_giorello/presentazione1.asp.
- [22] E. Giusti, *Analisi matematica 1*, Boringhieri, Torino, 1983.
- [23] R. Goldblatt, *Lectures on the Hyperreals: An Introduction to Nonstandard Analysis*, Springer, New York, 1998.

- [24] E. Hairer, G. Wanner, *Analysis by its History*, Springer-Verlag, New York Berlin Heidelberg, 2000.
- [25] C. Hermite, Sur la fonction exponentielle, *C. R. Académie des Sciences*, 77, 1873, 18–24, 74–79, 226–233, 285–293.
- [26] I.N. Herstein, *Algebra*, Editori Riuniti, Roma, 1982.
- [27] D. Hilbert, *Jahres. der Deut. Math.-Verein.*, 8, 1899, 180–184.
- [28] D. Hilbert, Mathematische Probleme, Vortrag, gehalten auf dem internationalen Mathematiker-Kongreß zu Paris 1900. URL: www.mathematik.uni-bielefeld.de/~kersten/hilbert/rede.html.
- [29] D. Hilbert, *Über den Zahlbegriff*. Jahresbericht der Deutschen Mathematiker-Vereinigung 8, pp. 180–184; trad. inglese: *On the concept of number*, in: W. Ewald, Ed., *From Kant to Hilbert: A source book in the foundations of mathematics*, Vol. 2, pp. 1089–1095, Oxford, Clarendon Press, 1966.
- [30] H. J. Keisler, *Elementi di analisi matematica*, Vol. I, Piccin Ed., Padova, 1982.
- [31] H. J. Keisler, *Elementary Calculus: An Approach Using Infinitesimals*, Prindle, Weber & Schmidt, II rev. ed., Boston, 1986. Ed. online (CCL). URL: www.math.wisc.edu/~keisler/keislercalc1.pdf.
- [32] J. Kennedy, *Completeness Theorems and the Separation of the First and Higher-Order Logic*, preprint, Università di Utrecht, gennaio 2008.
- [33] M. Kline, *Mathematical Thought From Ancient to Modern Times*, Oxford University Press, Oxford, 1972.
- [34] M. Kline, *Storia del pensiero matematico. I. Dall'Antichità al Settecento. II. Dal Settecento a oggi*, Einaudi, Torino, 1991.
- [35] G. Lakoff, R. E. Núñez, *Where Mathematics Comes From: How the Embodied Mind Brings Mathematics into Being*, Basic Books, 2000; trad. it.: O. Robutti, F. Ferrara, e C. Sabena, *Da dove viene la matematica. Come la mente embodied dà origine alla matematica*, Bollati Boringhieri, Torino, 2005. URL: www.cogsci.ucsd.edu/~nunez/web/.

- [36] F. Lindemann, Über die Zahl π , *Math. Ann.*, 20, 1882, 213–225.
- [37] G. Lolli, *Momenti di svolta nello sviluppo del pensiero matematico, Lezioni di Modena*, Modena, 2006.
URL: www2.dm.unito.it/paginepersonali/lolli/modena.htm.
- [38] G. Longo, The Mathematical Continuum, from Intuition to Logic, in : *Naturalizing Phenomenology: issues in contemporary Phenomenology and Cognitive Sciences*, J. Petitot, F. Varela, B. Pachoud, J–M. Roy, Eds., Stanford University Press, 1999.
- [39] C. Mangione, S. Bozzi, *Storia della logica. Da Boole ai nostri giorni*, Garzanti, Milano, 1993.
- [40] M. Maschietto, *L'enseignement de l'analyse au lycée : les debuts du jeu global / local dans l'environnement de calculatrices*, Tesi di dottorato, Université Paris 7 & Università di Torino, IREM Paris 7, 2002.
- [41] R. L. McCabe, Theodorus' Irrationality Proofs, *Mathematics Magazine*, Vol. 49, No. 4 (Sep., 1976), 201–203.
- [42] H. C. Méray, *Nouveau précis d'analyse infinitésimale*, F. Savy libraire–éditeur, Paris, 1872.
- [43] M. Michelotti Venè, C. Cervi, *et al.*, *Analisi non-standard: nozioni preliminari*, Quaderni del Dipartimento di Matematica, Università di Parma, 1996.
- [44] R. Nillsen, Normal Numbers without Measure Theory, *Amer. Math. Monthly*, Vol. 107, 7 (Aug.–Sep., 2000), 639–644.
- [45] R. E. Núñez, Do Real Numbers Really Move? Language, Thought, and Gesture: The Embodied Cognitive Foundations of Mathematics, in: F. Iida *et al.* (Eds.), *Embodied Artificial Intelligence*, LNAI 3139, Springer, 2004, 54–73.
- [46] R. E. Núñez, The Cognitive Science of Mathematics: Why Is It Relevant for Mathematics Education?, in: R.A. Lesh *et al.* (Eds.), *Foundations for the future in Mathematics Education*, Lawrence Erlbaum Associates, 2007, 127–154.

- [47] E. G. Omodeo, D. Cantone, A. Policriti, e J. T. Schwartz, A Computerized Referee, in: *Reasoning, Action and Interaction in AI Theories and Systems*, Lecture Notes in Computer Science, Springer, Berlin / Heidelberg, 2006.
- [48] Platone, *Theaetetus*, 147d.
- [49] R Development Core Team (2006). R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria. ISBN 3-900051-07-0. URL: www.R-project.org.
- [50] J. Richard, Les principes des mathématiques et le problème des ensembles, *Revue générale des sciences pures et appliquées*, 16, 1905, 541–543; trad. ingl.: in van Heijenoort, J., *From Frege to Gödel, A Source Book in Mathematical Logic*, Harvard Univ. Press, 1967, pp. 142–144.
- [51] A. Robinson, *Non-Standard Analysis*, North-Holland, Amsterdam, 1966.
- [52] C. Sabena, *Body and Signs: A Multimodal Semiotic Approach to Teaching-Learning Processes in Early Calculus*, Tesi di dottorato, Università di Torino, 2007.
- [53] M. Schiralli e N. Sinclair, A Constructive Response to ‘Where Mathematics Comes from’, *Educational Studies in Mathematics*, 52(1), 2003, 79–91.
- [54] O. Stolz, Zur Geometrie der Alten, insbesondere über ein Axiom des Archimedes, *Math. Ann.* 22, 1883, 504–520.
- [55] L. Talmy, Force dynamics in language and cognition, *Cognitive Science*, 12 (1), 1998, 49–100.
- [56] L. Talmy, Fictive motion in language and “ception”, in: P. Bloom, M. Peterson L. Nadel, M. Garrett (Eds.), *Language et space*, MIT Press, Cambridge, 1996.
- [57] A. Turing, On computable numbers, with an application to the Entscheidungsproblem, *Proc. London Math. Soc.*, Ser. 2, 42 (1936), 230–265. URL: www.abelard.org/turpap2/tp2-ie.asp.

- [58] A. N. Whitehead, B. Russell, *Principia Mathematica*, Cambridge University Press (2nd ed.), 1927.
- [59] P. B. Yale, Automorphisms of the Complex Numbers, *Mathematics Magazine*, 39(3), 1966, 135–141.
- [60] H. Weyl, *Das Kontinuum – Kritische Untersuchungen über die Grundlagen der Analysis*, Walter de Gruyter & Co., Berlin, 1918; trad. it.: *Il continuo - Indagini critiche sui fondamenti dell'analisi*, Bibliopolis, Napoli, 1977.

Fonti delle citazioni

- Cap. 1 DOXIADIS, A., *Zio Petros e la Congettura di Goldbach*, Bompiani, 2001.
- Cap. 2 THOM, R., *Mathématiques modernes, mathématiques de toujours*, in: *Pourquoi la Mathématique?*, Editore 10/18, Parigi, 1974.
- Cap. 3 CANTOR, G., *Grundlagen einer allgemeinen Mannigfaltigkeitslehre. Ein mathematisch-philosophischer Versuch in der Lehre des Unendlichen*, Teubner, Leipzig, 1883; trad. inglese in: EWALD, W. (ed.), *From Kant to Hilbert: A Source Book in the Foundations of Mathematics*, Oxford University Press, Oxford, 1996, pp. 639–920.
- Cap. 4 TITELMAN, G., *Random House Dictionary of Popular Proverbs and Sayings*, Random House USA Inc., New York, 1996.
- Cap. 5 LOCKSPEISER, E., *Music & Letters*, Vol. 15, No. 4 (Oct., 1934), p. 356.
- Cap. 6 GHERSI, I., *Matematica dilettevole e curiosa*, 5^a ed., Hoepli, 1988.
- Cap. 7 MOLIÈRE, *Le Bourgeois gentilhomme – . . .*, Gallimard, 2003; *Il borghese gentiluomo*, Rizzoli, 1996. (Atto II, Scena IV)
- Cap. 8 SHAKESPEARE, W., *Amleto*. Testo inglese a fronte, Garzanti Libri, 2008. (Atto I, Scena V)

Fonti delle immagini

<http://www-history.mcs.st-andrews.ac.uk/PictDisplay/>

Fig. 2.1 (a): [Hilbert.html](#)

Fig. 2.1 (b): [Archimedes.html](#)

Fig. 3.1 (b): [Meray.html](#)

Fig. 3.1 (b): [Cantor.html](#)

Fig. 4.1 (a): [Dedekind.html](#)

Fig. 4.1 (b): [Euclid.html](#)

Tutti gli URL citati nel testo sono stati visitati e verificati il 14 maggio 2008.

Indice analitico

A

allineamenti
 binari, 131
allineamento
 binario, 95, 98
 decimale, 95
anello quoziente, 47
assioma di scelta, 112
automorfismo, 40
Axiom der Vollständigkeit, 28, 42

C

campo, 13
 ordinato, 13
 archimedeo, 17, 21, 35
 completo, 20, 21, 35, 57, 82
 non-archimedeo, 17, 108, 119
classi
 contigue, 38, 68, 105
 separate, 38, 119
corpo, 13

D

denso, 19
derivata non-standard, 123
diagonale di Cantor, 131

E

elemento
 infinitesimo, 18, 108, 118
 infinito, 18

positivo, 14

separatore, 38, 119

estremo superiore, 7, 20

F

filtro, 109
 di Fréchet, 110
formula
 approssimazione lineare, 123

I

ideale massimale, 47, 114, 121
insieme
 numerabile, 128, 130
 trascurabile, 139
 triadico di Cantor, 141
iperreali, 108, 114
 finiti, 120
 infinitesimi, 118
 infiniti, 120

L

lemma di Zorn, 111, 113

N

numeri
 irrazionali, 134
 trascendenti, 134
 algebrici, 133
 calcolabili, 133
 non-calcolabili, 134

normali, 142

O

oggetto finale, 25

omomorfismo

ordinato, 22, 25

P

pacchetto $\mathbb{R}$, 4

paradosso

di Richard, 137

di Russell, 133

parte standard, 120

proprietà

di Archimede, 17, 38, 72

di Cantor, 36

di Dedekind, 36

Q

quadrato, 5

R

$\mathbb{R}$

definizione assiomatica, 27

modello

allineamenti di cifre, 104

classi contigue, 105

di Dedekind, 68

di Méray–Cantor, 47

$\mathbb{R}$ (pacchetto), 4

$^*\mathbb{R}$ (campo iperreale), 114

radice quadrata, 2

rappresentazione binaria

con *Mathematica*, 103

con $\mathbb{R}$, 99

S

segmento, 154

sezione

di Dedekind, 66, 95

di prima specie, 69

di seconda specie, 69

sottocampo razionale, 16

successione

convergente a 0, 46

di Cauchy

in $\mathbb{Q}$, 46

in $\mathbb{R}$, 53

in *AetnaNova*, 64

in Arzelà, Bolzano, 62

secondo Cantor, 59

T

teorema

della parte standard, 120

di Beigel, 7

di Bolzano–Weierstrass, 89

di Borel, 143

di Cantor, 88

di completezza metrica di $\mathbb{R}$, 89

di esistenza

dell'estremo superiore, 9, 57,

74, 83, 88

di Heine–Borel, 90

di Hilbert, 23

fondamentale

di Dedekind, 73, 83

di Méray–Cantor, 54

indecidibilità di Turing, 136

U

ultrafiltro

principale, 111

V

valore assoluto, 19

Glossario essenziale

Il Glossario richiama, per i non specialisti ed in modo discorsivo, i significati “intuitivi” dei principali concetti utilizzati nel testo.

ANELLO. Una struttura algebrica astratta all'interno della quale è possibile effettuare le operazioni di addizione, sottrazione e moltiplicazione, ma non sempre la divisione. Un esempio classico è l'anello $\mathbb{Z}$ dei *numeri interi*.

ASSIOMA DI ARCHIMEDE. Nasce nella teoria delle grandezze. Un sistema $\mathcal{G}$ di grandezze omogenee soddisfa l'Assioma di Archimede se, date due grandezze arbitrarie di $\mathcal{G}$, esiste un multiplo di una che supera l'altra. Analogamente, un campo ordinato $\mathbb{K}$ soddisfa l'Assioma di Archimede se, dati due elementi positivi arbitrari a, b di $\mathbb{K}$, esiste un multiplo $na = a + a + \dots + a$ (n volte) di a tale che $na > b$.

ASSIOMA DI COMPLETEZZA DI HILBERT (*Axiom der Vollständigkeit*). Introdotto da D. Hilbert nel 1900, questo assioma caratterizza il campo $\mathbb{R}$ dei numeri reali come il “più grande” campo ordinato in cui vale l'Assioma di Archimede.

CAMPO. Una struttura algebrica astratta nella quale è possibile effettuare le “quattro operazioni” di addizione, sottrazione, moltiplicazione e divisione per un elemento diverso da zero, in modo che siano verificate le usuali proprietà (associativa, commutativa, distributiva) delle omonime operazioni aritmetiche.

CAMPO ORDINATO. Un *campo* i cui elementi sono totalmente ordinati in modo “compatibile” con le operazioni.

COMPLETO. L'aggettivo *completo* riferito al sistema $\mathbb{R}$ dei numeri reali ha due significati diversi, ma equivalenti, quello di Dedekind, legato solo alla relazione d'ordine, e quello di Mèray–Cantor, che è legato al concetto di

distanza. Le due nozioni di completezza di $\mathbb{R}$ svolgono lo stesso ruolo, quello di garantire un teorema di esistenza del limite di opportune successioni, e sono equivalenti all'Assioma di completezza di Hilbert (*Axiom der Vollständigkeit*).

DEFINIZIONE ASSIOMATICA. La definizione assiomatica di un oggetto matematico non precisa la natura dell'oggetto, ma solo le regole alle quali l'oggetto deve sottostare. Ad esempio, la definizione assiomatica di probabilità (A. N. Kolmogorov, 1933) non dice cosa sia la probabilità, ma assegna le regole fondamentali del calcolo delle probabilità. Analogamente, la definizione assiomatica dei numeri reali (D. Hilbert, 1900) non ci dice cosa sia un numero reale, ma assegna le proprietà relative alle operazioni di addizione e moltiplicazione ed alla relazione d'ordine.

DOMINIO DI INTEGRITÀ. Un *anello* (non ridotto al solo zero) nel quale per l'operazione di moltiplicazione valgono la proprietà commutativa (ossia $ab = ba$) e la legge di annullamento del prodotto (ossia se $a \neq 0$ e $b \neq 0$ si ha $ab \neq 0$).

ESTREMO SUPERIORE. L'estremo superiore di un sottoinsieme E contenuto in un insieme totalmente ordinato X è il più piccolo elemento di X che è maggiore o uguale di ogni elemento di E . In $\mathbb{R}$ ogni insieme limitato E ha un estremo superiore. Invece, sia nel campo $\mathbb{Q}$ dei numeri razionali che nel campo ${}^*\mathbb{R}$ dei numeri iperreali esistono dei sottoinsiemi che, pur essendo limitati, sono privi di estremo superiore.

INFINITESIMO. Una grandezza a di un sistema $\mathcal{G}$ di grandezze omogenee è infinitesima se esiste una grandezza finita b di $\mathcal{G}$ che non può venir superata da alcun multiplo di a . In un campo ordinato $\mathbb{K}$ un elemento positivo a è infinitesimo se tutti i multipli $na = a + a + \dots + a$ (n volte) di a sono minori di $1_{\mathbb{K}}$. Un campo ordinato con elementi infinitesimi è il campo ${}^*\mathbb{R}$ dei *numeri iperreali*.

ISOMORFISMO. Un isomorfismo Φ di una struttura A in una struttura analoga B è una corrispondenza che associa ad ogni elemento a di A uno ed un solo elemento b di B (indicato con $\Phi(a)$), senza ometterne alcuno, in modo da “conservare” le strutture. Ad esempio, se le strutture includono un'operazione di addizione, un isomorfismo associa alla somma di due elementi di A la somma, eseguita in B , dei due elementi corrispondenti. Se le strutture includono una relazione d'ordine si ha $a' \prec_A a''$ se e solo se $\Phi(a') \prec_B \Phi(a'')$. Se esiste un isomorfismo di A in B si dice che A e B sono *isomorfi*; in tal caso, dal punto di vista della loro struttura, A e B sono indistinguibili.

NUMERI INTERI. I numeri interi sono i *numeri naturali* $0, 1, 2, \dots$ e gli *opposti* $-1, -2, -3, \dots$ dei numeri naturali positivi. I numeri interi formano un *anello*, indicato con $\mathbb{Z}$, perché Z è la lettera iniziale di “Zahl” (*numero*, in tedesco).

NUMERI IPERREALI. Il campo ${}^*\mathbb{R}$ dei *numeri iperreali* è un campo ordinato che contiene strettamente il campo $\mathbb{R}$ dei numeri reali e quindi contiene (per l'*Axiom der Vollständigkeit*) elementi infinitesimi. I numeri iperreali consentono di sviluppare rigorosamente l'analisi matematica in modo cosiddetto *non-standard* (A. Robinson, 1966), recuperando così il calcolo infinitesimale nella formulazione “algebrica” di Leibniz, nella quale la derivata è nient'altro che il quoziente di due infinitesimi mentre l'integrale è una somma (ancorché infinita) di infinitesimi.

NUMERI IRRAZIONALI. Sono gli oggetti che vengono “aggiunti” ai numeri razionali in modo da ottenere il campo ordinato *completo* $\mathbb{R}$. La loro natura non è univoca, ma quale che essa sia, il campo risultante $\mathbb{R}$ è sempre “lo stesso” (a meno di *isomorfismi*).

NUMERI RAZIONALI. Sono i numeri ottenibili come rapporto tra due numeri interi, il secondo dei quali diverso da 0. Ogni numero razionale quindi può essere espresso mediante una frazione. I numeri razionali formano un *campo*, indicato con $\mathbb{Q}$.

NUMERI REALI. Il campo $\mathbb{R}$ dei numeri reali è costituito dal campo $\mathbb{Q}$ dei numeri razionali (o uno a questo isomorfo), al quale vengono “aggiunti” degli oggetti matematici in modo da ottenere un campo ordinato *completo*. Gli oggetti aggiunti sono detti numeri irrazionali. Si noti che non c'è un criterio per stabilire se un singolo particolare oggetto matematico è o non è un numero reale: si può solo dire se un certo *insieme*, strutturato con una addizione, una moltiplicazione ed una relazione d'ordine totale, è o non è un campo ordinato completo, ossia se il dato insieme è o non è (isomorfo a) $\mathbb{R}$.

ORDINAMENTO. Un ordinamento, o *relazione d'ordine*, in un insieme non vuoto X è una relazione $\prec$ in X che soddisfa la proprietà riflessiva ($a \prec a$), antisimmetrica (da $a \prec b$ e $b \prec a$ segue $a = b$) e transitiva (da $a \prec b$ e $b \prec c$ segue $a \prec c$). L'ordinamento è *totale* se, quali che siano a, b in X , si ha $a \prec b$ oppure $b \prec a$. L'esempio classico di ordinamento totale è la relazione $\leq$ “minore o uguale” in un insieme X di numeri interi, o razionali, o reali.

RELAZIONE. Una relazione $\mathcal{R}$ in un insieme non vuoto X è formalmente un insieme $\mathcal{R}$ di coppie (a, b) di elementi di X . Se $(a, b) \in \mathcal{R}$ si scrive $a\mathcal{R}b$, e si dice che l'elemento a è in $\mathcal{R}$ -relazione con l'elemento b .