

TENTATIVI DI FONDARE LA MATEMATICA

Cosa sono i punti, le rette, i numeri ?

Giangiacoimo Gerla

**Versione ridotta e non aggiornata. La versione completa,
divisa in due volumi, può essere acquistata sul sito Ilmio-
libro**

<http://ilmiolibro.kataweb.it/vetrina.asp>

INDICE

Introduzione

CAPITOLO 1

LA MATEMATICA PRESSO I GRECI

1. La Scuola Pitagorica: tutto il mondo è aritmetica	1
2. Crisi della Scuola Pitagorica (ma gli interi non bastano)	5
3. Dimostrare e dimostrare per assurdo	11
4. Il continuo geometrico per evitare l'infinito attuale	13
5. Punti, linee e Platonismo	16
6. Gli elementi di Euclide	19
7. La teoria delle grandezze omogenee (al posto dei numeri reali) ...	25
8. La teoria delle proporzioni (al posto delle operazioni)	29
9. Misure, equiscomponibilità, equicompletabilità	32
10. L'equiscomponibilità è un metodo universale	36
11. Contro i matematici	42
Lettura: <i>Platone e la duplicazione del quadrato</i>	49

CAPITOLO 2

CRISI DELLA GEOMETRIA EUCLIDEA

1. Crisi del carattere assoluto della geometria	77
2. Modelli di geometrie non euclidee	81
3. Altre geometrie	84
4. Crisi dell'approccio sintetico: Cartesio	85
5. Calcolo dei segmenti	87
6. Il "Discorso sul Metodo"	89
7. La "costruzione" delle radici di una equazione	92
8. Aritmetizzazione della geometria: la sparizione delle figure	95
9. Intuizione geometrica e falsi teoremi euclidei	97

CAPITOLO 3

DEFINIRE I NUMERI

1. Un punto di partenza: terne di Peano	105
2. Principio di induzione	108
3. Definizione per ricorsione	110
4. Somma e prodotto in una terna di Peano.....	114
5. Definire una relazione d'ordine in una terna di Peano	115
6. Variazioni sul principio di induzione	119
7. L'anello degli interi relativi	122
8. Il campo dei razionali	125
9. I numeri reali tramite le sezioni	128

10. I numeri reali tramite le successioni di Cauchy	131
11. Un percorso diverso: essere quasi uguali	135
12. I razionali non-standard	138
Lettura: Zavattini, <i>Gara di matematica</i>	145

CAPITOLO 4

GLI INSIEMI: CREDERE NELL' INFINITO

1. Il prezzo dell'arimetizzazione: l'infinito attuale	147
2. Ma questi insiemi sono poi veramente una novità ?	150
3. I paradossi dell' infinito	153
4. Cantor, l'infinito e la Dottrina Cristiana	156
5. Confrontare le grandezze degli insiemi	158
6. Insiemi numerabili	161
7. Tentare di superare il numerabile	164
8. La potenza del continuo	168
9. Superare la potenza del continuo	173

Lettura: R. Rucker, <i>L'albergo di Hilbert</i>	193
--	-----

CAPITOLO 5

METODO ASSIOMATICO E STRUTTURALISMO

1. Paradossi e crisi della teoria degli insiemi	201
2. Russell, il paradosso del barbiere, Marx e le studentesse	205
3. Affrontare i paradossi: intuizionismo e metodo assiomatico	206
4. Un approccio assiomatico alla geometria	210
5. Un approccio assiomatico ai numeri reali	215
6. Assiomi per evitare i paradossi della teoria degli insiemi	218
7. La teoria di Zermelo-Frankel	220
8. Assioma della scelta	225
9. Dimostrare o confutare l'assioma della scelta	228
10. Ipotesi del continuo	229
11. Categoricità, consistenza, indipendenza, completezza	230
12. Tre diverse ideologie per il metodo assiomatico	233

CAPITOLO 6

LA MATEMATICA COME CALCOLO CON PAROLE

1. Hilbert contro l'infinito	243
2. L'infinito è solo una parola	246
3. Nuovi oggetti matematici: parole e linguaggi	248
4. Rappresentabilità, definibilità e numerabilità	252
5. Linguaggio ed apparato deduttivo per la logica formale	255
6. Ma si deve pur parlare di qualche cosa: l'interpretazione	260

7. Cosa è la verità	264
8. Teorema di completezza e teoremi limitativi	266

APPENDICE

NOZIONI BASE E VARIE

1. Coppie e prodotti cartesiani	281
2. Definizione (brutta) di n -pla	282
3. Relazioni di equivalenza e quozienti	284
4. Relazioni d'ordine e reticoli	286
5. Relazioni di buon ordine	288
6. Gruppi, anelli e campi	293
7. La nozione generale di struttura matematica	295
8. Sistemi di chiusura, operatori e punti fissi	299
9. Due teoremi di punto fisso per operatori	302
10. Come generare relazioni di ordine o di equivalenza	305

Indice analitico	321
-------------------------------	-----

Bibliografia	327
---------------------------	-----

INTRODUZIONE

Questo volume, che è una versione ridotta di un libro pubblicato sul sito “ilmiolibro”, si propone di introdurre il lettore alle problematiche relative ai fondamenti della matematica. Pertanto si occupa dei vari tentativi fatti da matematici e filosofi di fornire basi sicure alla matematica e di capirne i “principii primi” (ammesso che alla matematica si possa dare una base definitiva ed ammesso che per essa esistano principii primi). Nel libro si parte dalle idee che aveva in proposito Pitagora e si arriva fino a quelle di Hilbert. Tuttavia il libro non ha assolutamente intenzione di essere un libro di storia della matematica. Piuttosto la storia della matematica è utilizzata come filo conduttore di possibili percorsi di fondazione della matematica. Poiché gli appunti si riferiscono a studenti di un corso della laurea di Matematica, spesso nel libro si troveranno definizioni precise ed un po’ pedanti e, principalmente, dimostrazioni di teoremi. Tuttavia credo che il libro possa essere letto da tutti in quanto esiste l’antico diritto di chi legge di saltare pagine, dimostrazioni e parti noiose. Ovviamente i miei studenti sono esclusi da tale diritto.

Naturalmente si pone la questione:

esiste l’esigenza di un tale tipo di riflessione sulla matematica ?

Dopotutto sembra non esistere niente di più semplice e sicuro di nozioni come quelle di numero, punto, retta. Purtroppo semplicità e sicurezza sono illusioni come mostrano i vari paradossi emersi nel corso dell’evoluzione della matematica. Il fatto è che siamo tanto abituati a manipolare i concetti matematici che tendiamo a confondere la familiarità che abbiamo acquisito con la conoscenza di tali concetti. Un po’ avviene come per il nostro giornalista o salumiere che pensiamo di conoscere solo perché sono venti anni che facciamo acquisti da loro (ma poi non sappiamo nemmeno dove abitano o se sono sposati o no).

Proviamo però ad essere meno superficiali ed a porci domande del tipo:

- che cosa sono i numeri?
- che cosa è un punto, una retta?
- i numeri, i punti le rette sono invenzioni dell’uomo, di un dio oppure esistono in natura?
- che cosa è la matematica?
- i risultati della matematica sono sicuri? e, se sono sicuri, perché lo sono?

-
- esiste l'infinito di cui spesso parla la matematica ?
 - perché la matematica, che non sembra avere a che fare con l'esperienza, è utile per le scienze empiriche e per le applicazioni ?

(ho detto meno superficiali !!! chi avesse letto troppo rapidamente deve tornare indietro e riflettere su ciascuna domanda per almeno un minuto per cercare una possibile risposta). Ci accorgiamo allora che tali questioni sono molto più problematiche di quanto appare a prima vista. Inoltre, se ci si guarda un po' in giro, ci si accorge che persone diverse, specialmente se appartenenti ad epoche diverse, hanno dato e danno risposte diverse. Questo fatto si esprime dicendo che:

sono esistite ed esistono diverse "filosofie" della matematica.

Insomma il libro vuole introdurre chi legge a tali questioni rimandando a libri più "solidi" nel caso si volesse approfondire.

Chiudo questa introduzione dichiarando che sono cosciente che nel libro esistono molte lacune. Ad esempio avrei dovuto trattare la teoria delle categorie ed in particolare la teoria dei topos. Inoltre lo spazio dedicato alla trattazione dell'intuizionismo non corrisponde certo alla importanza ed all'originalità di questa filosofia della matematica. Infine esistono settori della matematica, come la probabilità, che pur presentando interessantissimi problemi fondazionali, non vengono trattati. Queste mancanze derivano dalla mia pigrizia e non dal misconoscimento di questi argomenti.

Salerno 2010

P.S.

Il mio indirizzo è gerla@unisa.it e ricevo volentieri commenti, segnalazioni di errori o richieste di chiarimenti.

CAPITOLO 1

LA MATEMATICA PRESSO I GRECI

E' indegno del nome di uomo chi ignora il fatto che la diagonale di un quadrato è incommensurabile con il suo lato. Platone (429-347 a.C.)

1. La scuola Pitagorica: tutto il mondo è aritmetica

Il primo organico tentativo di dare una fondazione alla matematica (ed all'intera conoscenza scientifica) fu probabilmente quello della scuola pitagorica il cui assunto di partenza era che:

alla base di tutto è il numero intero.

La scuola pitagorica era una setta mistico-religiosa che si sviluppò in Grecia e in Italia (Crotone) tra il 570 ed il 500 a.C. attorno ad un mitico personaggio chiamato Pitagora. Le idee di tale scuola sono di fondamentale importanza per la storia della cultura occidentale perché da esse inizierà quel processo che trasformerà la scienza pre-ellenica, che consisteva in una disarticolata raccolta di risultati dettati dall'esperienza, in una scienza razionale.¹ Dei pitagorici ne parla Aristotele al modo seguente, dove si deve tenere conto che allora per "numero" si intendeva "numero intero positivo".

Tra i primi filosofi, ..., furono i cosiddetti Pitagorici, i quali, applicatisi alle scienze matematiche, le fecero per i primi progredire; cresciuti poi nello studio di esse, vennero nell'opinione che i loro principi fossero i principi di tutti gli esseri... Pensarono che gli elementi dei numeri fossero gli elementi di

¹Non bisogna avere una immagine dei pitagorici come scienziati campioni di razionalismo. Il carattere mistico di questa scuola era fortissimo, siamo in presenza di una vera e propria setta religiosa (e politica) che credeva, tra le altre cose, che le anime dei morti si reincarnassero negli animali. Anche le "regole" di tale setta ci appaiono notevolmente bizzarre. Ad esempio ecco alcuni comandamenti:

- non toccare un gallo bianco
- non addentare una pagnotta intera
- non guardare uno specchio accanto ad un lume.

Ma queste che ci appaiono come stranezze non tolgono ai pitagorici il merito di costituire il punto di inizio della moderna cultura scientifica.

tutte le cose, e che l'universo intero fosse armonia e numero (Aristotele, Metafisica).

Si deve tenere conto che in quel periodo era forte il desiderio di trovare i "principi ultimi" e che questi venivano cercati negli elementi naturali come l'aria, l'acqua o il fuoco. Forse però per capire meglio il pensiero dei Pitagorici conviene vedere cosa dice uno di loro, Filolao.

“Nessuna menzogna accolgono in sé la natura del numero e l'armonia: non è cosa loro la menzogna. La menzogna e l'invidia partecipano della natura dell'illimitato, dell'intellegibile e dell'irrazionale. Nel numero non penetra menzogna, perché la menzogna è avversa e nemica della natura, così come la verità è connaturata e propria alla specie dei numeri . . . ”

“...Nulla sarebbe comprensibile, né le cose in sé né le loro relazioni, se non ci fossero il numero e la sua sostanza.”

“Tutte le cose che si conoscono hanno numero: senza il numero non sarebbe possibile pensare né conoscere alcunché.”

Un ruolo talmente centrale dato al numero potrebbe anche dipendere dalla scoperta di un forte collegamento tra rapporti numerici ed “armonia” in campo musicale. Infatti viene attribuita ai Pitagorici la scoperta di una scala armonica che viene detta, appunto, *scala pitagorica*. Consideriamo delle corde tese di varia lunghezza ed esaminiamo i suoni che vengono emessi pizzicandone due contemporaneamente. Ci si accorge che a volte si hanno effetti gradevoli ed a volte sgradevoli. E' possibile studiare quale sia il rapporto tra le lunghezze delle due corde ed il fenomeno della “gradevolezza” o, per essere più specifici, della “consonanza”. Ora la prima scoperta che viene da fare è che se una corda è il doppio dell'altra si ha una fortissima consonanza. In questo caso noi diciamo che i due suoni differiscono per una ottava. Se indichiamo con A la lunghezza della prima corda e con B quella della seconda allora $B = (1/2) \cdot A$ o, se si vuole, $A : B = 2 : 1$. Un altro suono gradevole si ottiene facendo vibrare, insieme ad A una corda la cui lunghezza C sia i due terzi di A cioè $C = (2/3) \cdot A$. Ne segue che $A : C = 3 : 2$. Infine ci si accorge che un suono gradevole si ottiene dai suoni prodotti dalle corde C e B che risultano essere nel rapporto $C : B = 4 : 3$. Abbiamo quindi che

le tre consonanze principali, (che prendono il nome di ottava quinta e quarta), corrispondono ai rapporti 2:1; 3:2 e 4:3. E' una sorprendente corrispondenza tra suoni e numeri che suggerisce fortemente l'idea per cui "tutto è numero".

Da quel "*tutte le cose che si conoscono hanno un numero*" scaturiva poi il convincimento circa la struttura granulare e discreta delle figure geometriche e, più in generale, del mondo fisico. Ciò comportava, ad esempio, una concezione del segmento come insieme finito di punti-unità, punti che venivano intesi come veri e propri corpi con una determinata grandezza. Infatti era solo in base a tale ipotesi che i numeri interi potevano essere lo strumento perfettamente adeguato alla descrizione della realtà, anzi, in un certo senso, venivano a coincidere con la realtà stessa. In tale modo la geometria non si considerava distinta dall'aritmetica e, in un certo senso, l'aritmetica assumeva una forma geometrica. Dei numeri infatti si dava una rappresentazione geometrica o, se si vuole, fisica, tramite una opportuna configurazione di punti-sassolino. Ad esempio si chiamavano *triangolari* i numeri che si potevano ottenere disponendo i punti in triangoli.

```

      O      O      O      O
        O O    O O    O O
          O O O  O O O
            O O O O
              O O O O

```

Si chiamavano invece *quadrati* i numeri corrispondenti a gruppi di sassolini disposti in quadrato

```

      O      O O      O O O      O O O O
        O O    O O O    O O O O    O O O O
          O O O  O O O O  O O O O O
            O O O O

```

I numeri quadrati corrispondono ai numeri che sono quadrati perfetti. Ancora, si chiamavano *rettangolari* i numeri corrispondenti a gruppi di sassolini disposti in un rettangolo (che non si riduca ad una striscia di sassolini). In questo caso sono rettangolari tutti e soli i numeri che non sono numeri primi.

```

      O O      O O O      O O O      O O O O
      O O      O O O      O O O      O O O O
                O O O      O O O O

```

L'interpretazione dei numeri come particolari disposizioni di sassolini consente di sviluppare una interessante aritmetica. Ad esempio, è immediato che ogni triangolo si ottiene dal precedente aggiungendo un fila di sassolini. Pertanto se $t(n)$ è il numero dei sassolini dell'ennesimo triangolo abbiamo che la funzione t è definibile tramite le equazioni

$$t(1) = 1 : t(n) = t(n-1) + n.$$

Ne segue che sono triangolari tutti i numeri della serie 1, 3, 6, 10, ... di termine generale n , cioè tutti i numeri del tipo $1+2+3+\dots+n$. Per quanto riguarda i numeri quadrati, è immediato vedere che ogni quadrato si ottiene dal precedente aggiungendo due lati (con un sassolino in comune). Ne segue che, se $q(n)$ è il numero dei sassolini dell'ennesimo quadrato, la funzione q si definisce tramite le equazioni

$$q(1) = 1 : q(n) = q(n-1) + 2n - 1.$$

Pertanto i quadrati perfetti si ottengono come elementi della serie 1, 1+3, 1+3+5, ..., $1+3+\dots+2n-1$, cioè la serie di termine generale $2n-1$. Si osservi che sia la funzione t che la funzione q sono state definite "per ricorsione". In proposito si veda nel prossimo capitolo i paragrafi su induzione e ricorsione.

L'importanza della svolta impressa dalla scuola pitagorica non si limita alla sola matematica poiché la fede nella potenza regolarizzatrice del numero intero, il procedimento di astrazione, l'uso delle dimostrazioni nel procedere scientifico, rappresentano il nascere dell'aspetto fondamentale della cultura occidentale: il convincimento che il mondo sia comprensibile non attraverso l'ascesi mistica, la contemplazione, come viene ritenuto dalle culture orientali, ma attraverso l'attività raziocinante. Con Pitagora ha inizio un processo di idealizzazione e razionalizzazione di tutte le forme di conoscenza che dominerà perfino la nostra cultura religiosa. Afferma ad esempio Bertrand Russell in *"Storia della filosofia occidentale"* :

"La religione razionalistica, al contrario di quella apocalittica, è stata da Pitagora in poi (ed in particolare da Platone in poi) completamente dominata dalla matematica e dal metodo matematico. La combinazione di matematica e di teologia, che cominciò con Pitagora, caratterizzò la filosofia religiosa in Grecia, nel Medioevo e nell'era moderna fino a Kant. L'orfismo precedente a Pitagora era analogo alle misteriose religioni asiatiche. Ma, in Platone, Sant'Agostino, Tommaso d'Aquino, Cartesio, Spinoza e Leibniz, vi è un intimo intrecciarsi

di religione e di ragionamento, di aspirazione morale e di ammirazione logica per ciò che è eterno, il quale viene da Pitagora e distingue la teologia intellettualizzata dell'Europa dal più diretto misticismo asiatico.”

2. Crisi della scuola pitagorica (ma gli interi non bastano).

Lo sapevate? Il quadrato costruito sull'ipotenusa è il doppio di quello sui cateti . . . ma la qualità è scadente e dopo un anno lo butti! È così! È capitato a mia sorella! Fidatevi!

-Vulvia (Corrado Guzzanti), Il caso Scrafoglia, 2002

Le concezioni dei pitagorici furono però ben presto messe in crisi dalla scoperta della esistenza di grandezze geometriche "incommensurabili". Questa scoperta è conseguenza del teorema che va proprio sotto il nome di "Teorema di Pitagora".

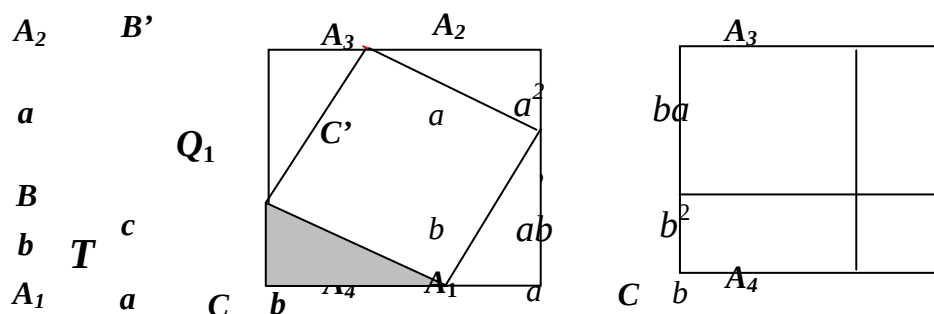
Teorema 2.1. (Teorema di Pitagora) Dato un triangolo rettangolo, se si considera l'unione dei due quadrati costruiti sui cateti otteniamo una figura che ha la stessa estensione del quadrato costruito sull'ipotenusa.²

*Dim.*³ Indichiamo con T il triangolo rettangolo e supponiamo che i cateti misurino a e b e l'ipotenusa c . Costruiamo un quadrato Q

² Attualmente tale teorema si enuncia dicendo che: "la somma dei quadrati delle misure dei cateti è uguale al quadrato della misura dell'ipotenusa". Tuttavia si deve tenere presente che per i greci non esisteva una "misura" di una figura geometrica in quanto misurare significa assegnare un numero reale a qualche cosa e nella matematica greca non esisteva una nozione di numero reale. Come vedremo, essi avevano invece la nozione di uguaglianza dell'estensione di due figure, quella di equiscomponibilità e quella di proporzionalità di grandezze omogenee con cui riuscivano ad esprimere molti teoremi sulle lunghezze, le aree ed i volumi.

³ In realtà quella che segue non è una dimostrazione nel senso rigoroso del termine e non lo può essere visto che non abbiamo ancora elencato gli assiomi della geometria di Euclide. Piuttosto si prova una proprietà non molto evidente "il teorema di Pitagora" a partire da altre proprietà

con lati uguali ad $a+b$. Detti A_1, A_2, A_3, A_4 i vertici di Q , tracciamo sui lati i quattro punti C, C', B, B' in modo che $A_1C = A_4C' = A_3B' = A_2B = a$ mentre $CA_4 = C'A_3 = B'A_2 = BA_1 = b$. In tale modo si individuano quattro triangoli rettangoli uguali (avendo cateti uguali per costruzione). Inoltre tali punti individuano un quadrilatero $Q_1 = BCC'B'$. Tale quadrilatero ha lati uguali in quanto coincidenti con le ipotenuse dei triangoli. Inoltre gli angoli sono retti.



Ad esempio, l'angolo in C è retto in quanto è uguale ad un angolo piatto meno la somma dei due angoli non retti di T . D'altra parte la somma dei due angoli non retti di un triangolo rettangolo è un angolo retto. Pertanto Q_1 è un quadrato.

Osserviamo ora che l'area di Q può essere calcolata come il quadrato del lato cioè come $(a+b)^2$ e quindi, per la formula del quadrato di un binomio, come a^2+b^2+2ab . D'altra parte l'area di Q è anche uguale al quadrato piccolo Q_1 più i quattro triangoli, cioè è uguale a $c^2+(4ab)/2$. Dall'uguaglianza

$$a^2+b^2+2ab = c^2+2ab$$

si ricava che $a^2+b^2=c^2$. Da notare che la formula del quadrato del binomio ammette una semplice dimostrazione geometrica che è completamente illustrata dalla figura tracciata sopra a destra in cui si mostra come l'area del quadrato si possa calcolare in due modi. Uno è considerare il quadrato di $a+b$, l'altro è effettuare la

che ci appaiono più evidenti. Ad esempio il primo passo della dimostrazione consiste nella costruzione di un quadrato di un dato lato. E' possibile sempre costruire un tale quadrato? La questione non è tanto semplice ed ha a che fare, tra l'altro, con l'assioma delle parallele. Nel corso di questo capitolo considereremo spesso "dimostrazioni" di questo tipo.

somma dei quadrati di estensione a^2 e b^2 più due rettangoli di estensione ab . $\square$

E' possibile anche porsi il seguente problema:

“un triangolo che verifica il teorema di Pitagora è necessariamente rettangolo ?

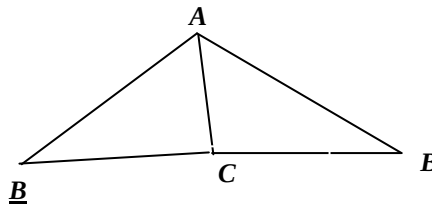
Ad esempio, supponiamo di volere formare un triangolo con tre segmenti di lunghezza 3, 4, 5. Tenendo conto del fatto che $3^2 + 4^2 = 5^2$ è possibile affermare che il triangolo è rettangolo ? Il seguente teorema mostra che la risposta è positiva.

Teorema 2.2. (Teorema inverso di Pitagora) Ogni triangolo i cui lati verificano il teorema di Pitagora è rettangolo.

Dim. Sia ABC un triangolo tale che $AB^2 = AC^2 + BC^2$ e costruiamo un segmento $\underline{BC}$ perpendicolare ad AC e di lunghezza uguale a CB .

Allora i due triangoli $\underline{ACB}$ e ACB hanno il lato AC in comune ed i lati $\underline{BC}$ e CB uguali per costruzione. Inoltre, essendo $\underline{ACB}$ rettangolo in C , $AB^2 = AC^2 + \underline{BC}^2 = AC^2 + CB^2 = AB^2$.

Pertanto i due triangoli, avendo i tre lati uguali, sono uguali. Da ciò segue che l'angolo ACB è uguale all'angolo retto $\underline{ACB}$ ed è quindi retto. $\square$



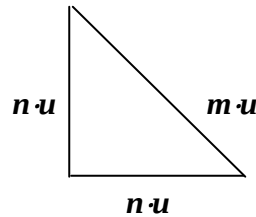
Dal teorema di Pitagora segue un sorprendente paradosso⁴. Dato un quadrato, comunque si scelga un segmento u come unità di misura non possiamo misurare il suo lato e la sua diagonale tramite due interi.

Teorema 2.3. (Paradosso dell' incommensurabilità tramite gli interi) Il lato e la diagonale di un quadrato sono incommensurabili.

⁴ Letteralmente “paradosso” significa “contro l'opinione comune” e tale espressione non andrebbe confusa con “antinomia” o “contraddizione” che invece significano che si è riusciti a provare una affermazione ed anche la negata di tale affermazione. Tuttavia spesso si confondono le due cose anche perchè, se si aggiunge l' “opinione comune” come assioma, allora il paradosso diviene una contraddizione.

surabili, cioè comunque si fissi un segmento u come unità di misura, il lato e la diagonale non sono uguali a multipli di u .

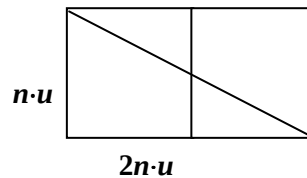
Dim. La dimostrazione si basa sul fatto che in un quadrato perfetto q ogni fattore primo è presente un numero pari di volte. Infatti se $q = n^2$ e se n ammette la scomposizione $n = p_1^{n(1)} \dots p_s^{n(s)}$ con $p_1, \dots, p_s$ primi allora $q = n^2 = p_1^{2 \cdot n(1)} \dots p_s^{2 \cdot n(s)}$. Ad esempio se $q = 12^2 = 144$ allora essendo $12 = 2^2 \cdot 3$, sarà $q = 12^2 = 2^4 \cdot 3^2$ e quindi q contiene il numero primo 2 quattro volte ed il numero primo 3 due volte. Supponiamo ora per assurdo che esista un segmento u (unità di misura) tale che sia il lato che la diagonale siano multipli, secondo gli interi n ed m , di u . In altri termini,



supponiamo che il lato e la diagonale siano misurati tramite due numeri interi n ed m . Allora i quadrati costruiti sui cateti sono costituiti da n^2 quadratini unitari mentre il quadrato costruito sull'ipotenusa è costituito da m^2 quadratini unitari. Ne segue che per il teorema di Pitagora $n^2 + n^2 = m^2$ e quindi $2n^2 = m^2$. Ma ciò è assurdo poiché, posto $c = 2n^2 = m^2$,

- dall'equazione $c = 2n^2$ si deduce che il fattore 2 è presente in c un numero dispari di volte
- dall'equazione $c = m^2$ si deduce che il fattore 2 compare in c un numero pari di volte. □

Problema: Consideriamo un rettangolo con un lato uguale al doppio dell'altro. La diagonale è commensurabile con il lato minore?



Abbiamo visto che i numeri interi non costituiscono uno strumento sufficiente per effettuare misure geometriche. La cosa non

migliora se si coinvolgono i numeri razionali. Infatti vale anche il seguente teorema.

Teorema 2.4. (Paradosso dell' incommensurabilità tramite i razionali) Dato un quadrato, per quanto si scelga piccolo un segmento come unità di misura⁵, le misure del lato e della diagonale rispetto tale unità non possono essere espresse da numeri razionali.⁶

Dim. Supponiamo per assurdo che esista un segmento unitario u tale che le lunghezze del lato e della diagonale siano esprimibili tramite due razionali. Allora, detti p/q ed r/s tali razionali avremmo, sempre per il teorema di Pitagora, che $2(p/q)^2 = (r/s)^2$ e quindi che

$$2(p \cdot s)^2 = (r \cdot q)^2.$$

Posto $n = p \cdot s$ ed $m = r \cdot q$ abbiamo che $2 \cdot n^2 = m^2$. Ma abbiamo già visto che una tale equazione conduce all'assurdo. $\square$

In definitiva, in termini attuali, diremmo che perfino la misurazione di grandezze legate alla figura geometrica più elementare, il quadrato, comporta necessariamente il coinvolgimento dei numeri irrazionali. Ora, se si tiene conto del fatto che per gli antichi greci i soli numeri esistenti erano gli interi positivi, ciò significava per la cultura del tempo che:

vi sono più cose in geometria (e quindi nel mondo fisico) di quanto i numeri siano capaci di esprimere !

D'altra parte, poiché la razionalità veniva identificata con il numero (cioè la possibilità di esprimere tramite numeri e rapporti di

⁵ L'espressione "per quanto si scelga piccolo" ovviamente non ha senso in matematica. Infatti la nozione vaga di "piccolo" non ha carattere matematico o logico. Il matematico pignolo pertanto è autorizzato a sostituirla con l'espressione "per ogni". La stessa situazione si presenta ad esempio nella definizione di limite dove si usano addirittura espressioni del tipo "preso ε piccolo a piacere ...", dove alla vaghezza del "piccolo" si aggiunge quella ancora maggiore di "a piacere". D'altra parte chiunque abbia esperienza didattica "sente" che queste espressioni sono utili.

⁶ Questa volta abbiamo usato espressioni moderne del tipo "la misura del lato l rispetto ad un segmento u è uguale al numero razionale p/q ". Per esprimere tale proprietà i greci avrebbero detto "il lato l ed il segmento u sono in proporzione con p e q ". Si veda in proposito la teoria delle proporzioni.

numeri i fenomeni del mondo), una tale conclusione comportava la messa in discussione della stessa possibilità, da parte dell'uomo, di pervenire alla conoscenza.⁷ Un altro fatto importante è il seguente:

alla luce della scoperta di grandezze incommensurabili come il lato e la diagonale del quadrato la stessa concezione dei segmenti ed in generale di ogni figura geometrica come somma finita di punti-atomo, che era propria della scuola pitagorica, non poteva più reggere.

Se infatti i segmenti fossero costituiti da un numero finito di particelle indivisibili, tutte di lunghezza u , allora ovviamente il lato e la diagonale del quadrato avrebbero lunghezza multiplo di u e sarebbero quindi commensurabili. D'altra parte si tenga presente che una tale concezione coincide proprio con quella della fisica moderna che sembra pertanto essere in contrasto con la geometria, o, per meglio dire, con il teorema di Pitagora.

Concludiamo questo paragrafo osservando che:

se con la scuola pitagorica abbiamo il primo tentativo di fondazione generale della matematica, con la scoperta delle grandezze incommensurabili siamo in presenza della prima “crisi dei fondamenti” della matematica.

3. Dimostrare e dimostrare per assurdo

Con il Teorema 2.1 ed il Teorema 2.3 abbiamo visto due tra i più antichi esempi di “dimostrazione”. Tali dimostrazioni sono di natura molto diversa. Nel Teorema 2.1 si prova qualcosa in positivo, precisamente che vale una certa equazione e la dimostrazione consiste nel mettere in evidenza che valgono una serie di uguaglianze. La dimostrazione della incommensurabilità del lato

⁷ Come fosse sconvolgente per i greci una tale scoperta possiamo vederlo attraverso questo passo attribuito al filosofo Proclo Diadoco.

“E' fama che colui il quale per primo rese di pubblico dominio la teoria degli irrazionali sia perito in un naufragio, e ciò perché l'insprimibile e l'inimmaginabile sarebbero dovuti rimanere sempre celati. Perciò il colpevole, che fortuitamente toccò e rivelò questo aspetto delle cose viventi, fu trasportato al suo luogo di origine e là viene in perpetuo flagellato dalle onde.”

e della diagonale di un quadrato è invece diversa e costituisce uno dei primi esempi di “*dimostrazione per assurdo*”. In questo caso si esprime una cosa in negativo, cioè che non vale una affermazione. Stante l'importanza di tale tipo di dimostrazione, esaminiamo da vicino la sua struttura.

Ricordiamo che chiamiamo *contraddizione* l'affermazione di un fatto e la contemporanea negazione di tale fatto. Indicando con B una asserzione, con $\neg B$ la sua negata, con $\wedge$ la congiunzione logica “e”, allora una contraddizione ha una forma del tipo $B \wedge (\neg B)$ che si legge “ B e non B ”. Naturalmente una contraddizione non può essere un'asserzione vera. Inoltre, poiché da asserzioni vere si deducono ancora asserzioni vere:

se da una asserzione A segue una contraddizione,

allora A non può essere vera,

pertanto : $\neg A$ è vera.

In definitiva la struttura di una dimostrazione per assurdo di una proposizione $\neg A$ è questa:

1. si suppone A
2. da tale ipotesi si ricava una contraddizione $B \wedge \neg B$,
3. si conclude che, non potendo valere A , vale $\neg A$.

Ad esempio, nel caso della dimostrazione di incommensurabilità del lato e della diagonale del quadrato,

- si suppone la commensurabilità,
- da tale ipotesi si ricava sia l'affermazione $B =$ “in c il fattore 2 è presente un numero pari di volte” sia l'affermazione $\neg B =$ “in c il fattore 2 è presente un numero dispari di volte”.
- si conclude che non sussiste la commensurabilità.

Naturalmente le dimostrazioni per assurdo possono essere utilizzate anche per provare qualche cosa in positivo. Possiamo infatti affermare il seguente altro principio:

se da una asserzione $\neg A$ segue una contraddizione,

allora $\neg A$ non può essere vera

pertanto: A è vera.

Pertanto, dovendo dimostrare A ,

1. si suppone $\neg A$
2. da tale ipotesi si ricava una contraddizione $B \wedge \neg B$,
3. si conclude che, non potendo valere $\neg A$, vale A .

Un'altra bella dimostrazione per assurdo è la dimostrazione per cui i numeri primi sono infiniti.

Teorema 3.1. L'insieme P dei numeri primi è infinito.

Dim. Detto A l'enunciato " P è infinito" neghiamo, cioè supponiamo per assurdo che l'insieme P dei numeri primi sia finito. Posto $P = \{p_1, \dots, p_n\}$, sia $q = p_1 \cdot \dots \cdot p_n + 1$: vogliamo dimostrare che q è primo. Infatti se per assurdo q non fosse primo ammetterebbe un divisore primo $p \neq 1$. Poiché abbiamo supposto che $p_1, \dots, p_n$ sono tutti i possibili numeri primi, p deve coincidere con un opportuno p_i . Allora p_i , dividendo sia q che $p_1 \cdot \dots \cdot p_n$ dovrà dividere anche $q - p_1 \cdot \dots \cdot p_n = 1$, cosa questa assurda. L'assurdo a cui siamo pervenuti prova che q è primo. D'altra parte q , essendo maggiore dei numeri $p_1, \dots, p_n$, non appartiene a $\{p_1, \dots, p_n\}$, in contrasto con l'ipotesi che $\{p_1, \dots, p_n\}$ è l'insieme di tutti i numeri primi.

Per provare questo teorema abbiamo dovuto utilizzare due volte il metodo di dimostrazione per assurdo. Possiamo eliminare una di queste utilizzazioni e riformulare il teorema in modo più "costruttivo".

Teorema 3.2. Dato un insieme finito $\{p_1, \dots, p_n\}$ di numeri primi esiste un numero primo p che non appartiene a $\{p_1, \dots, p_n\}$. Ne segue che l'insieme dei numeri primi è infinito.

Dim. Consideriamo il numero $q = p_1 \cdot \dots \cdot p_n + 1$ e sia p un divisore di q diverso da 1 (che potrebbe coincidere con q se q fosse primo). Se per assurdo $p \in \{p_1, \dots, p_n\}$ allora p , dividendo sia q che $p_1 \cdot \dots \cdot p_n$, dividerebbe anche la differenza $q - p_1 \cdot \dots \cdot p_n = 1$ e ciò è assurdo. $\square$

Ad esempio, dato l'insieme $\{2, 3, 7\}$ di numeri primi, il numero $2 \cdot 3 \cdot 7 + 1 = 43$ è primo. Dato l'insieme $\{3, 7, 11\}$ di numeri primi, il numero $3 \cdot 7 \cdot 11 + 1 = 232$ non è primo ma è divisibile per il numero primo 2 che è diverso da 3, 7, 11.

Esempio. Quando si studiano le equazioni si dimostrano spesso "teoremi" per assurdo. Ad esempio proviamo il seguente "teorema":

$A = \text{“L'equazione } x^2 + 2 \cdot (1 + x^2) = 3x^2 - 2 \text{ non ammette soluzioni.”}$
 Se neghiamo A dobbiamo accettare l'esistenza di un elemento r tale che

$$r^2 + 2 \cdot (1 + r^2) = 3r^2 - 2$$

da ciò seguirebbe che

$$r^2 + 2 + 2 \cdot r^2 = 3r^2 - 2 \text{ e quindi}$$

$$3 \cdot r^2 + 2 = 3r^2 - 2 \text{ e quindi, semplificando } 3 \cdot r^2$$

$$2 + 2 = 0 \text{ cioè}$$

$$4 = 0.$$

L'assurdo a cui siamo pervenuti ci assicura che vale A .⁸

4. Il continuo geometrico per evitare l'infinito attuale

Il fatto che i numeri interi si rivelassero uno strumento inadeguato a definire le grandezze geometriche poteva, da un punto di vista tecnico, essere risolto (almeno) in due modi diversi.

1. Si poteva ampliare il concetto di numero,
2. Si poteva decidere che la geometria non è riconducibile all'algebra, cioè alla nozione di numero.

I Greci seguirono la seconda via. Il primo punto di vista sarà invece assunto, come vedremo nel seguito, dalla matematica moderna con la definizione dei numeri reali e la successiva loro utilizzazione per la costruzione del continuo geometrico (la famosa geometria analitica). Poiché i reali si definiscono a partire dai razionali, e questi a partire dagli interi, l'attuale punto di vista sembra il naturale sviluppo di quello della scuola pitagorica. D'altra parte i Greci non potevano definire i reali perché non è possibile definire i numeri reali a partire dagli interi senza coinvolgere la nozione di infinito attuale. Basta osservare che, come faremo in seguito, un numero reale si definisce come un insieme attualmente infinito di razionali (si veda il metodo delle sezioni) oppure come una successione attualmente infinita di cifre decimali.

Ma gli antichi greci rifiutavano l'infinito attuale

⁸ In realtà non siamo pervenuti ad un assurdo poiché l'equazione $4 = 0$ è da considerare assurda solo se siamo nell'ambito di teorie in cui sia dimostrabile $\neg(4=0)$ come avviene nella teoria dei numeri reali. Se invece siamo, ad esempio, nell'anello degli interi modulo 4 la conclusione non è affatto assurda.

“... che il numero è infinito in potenza, ma non in atto ... questo nostro discorso non intende sopprimere per nulla le ricerche dei matematici per il fatto che esso esclude che l'infinito per accrescimento sia tale da poter essere percorso in atto. In realtà essi stessi (i matematici), allo stato presente, non sentono il bisogno dell'infinito (e in realtà non se ne servono) ma soltanto di una quantità grande quanto essi vogliono, ma pur sempre finita ... “ (Aristotele).

Ancora, Aristotele (Fisica III) afferma che l'infinito è tale

“... che si può prendere sempre qualcosa di nuovo (in esso), e ciò che si prende è sempre finito ma sempre diverso. Sicché non bisogna prendere l'infinito come un singolo essere, per esempio un uomo o una cosa, ma nel senso in cui si parla di una giornata o di una lotta, il cui modo d'essere non è una sostanza ma un processo e che, se pure è finito, è incessantemente diverso.”

Esistono mille esempi in cui si manifesta questo rifiuto dell'infinito da parte dei Greci. Ad esempio, il teorema da noi dimostrato circa l'esistenza di infiniti numeri primi in realtà veniva enunciato dai Greci al modo seguente in cui non viene coinvolta la nozione di insieme infinito.

Per ogni primo p esiste un primo q maggiore di p .

Inoltre in geometria non si concepiva la retta intesa come qualche cosa di illimitato ma ci si riferiva solo ai segmenti. Naturalmente veniva accettato che un segmento si potesse prolungare a piacere.

Il motivo per cui i Greci avessero tanta repulsione per l'infinito attuale è probabilmente di natura strettamente filosofica ed è ben illustrato dal passo dei pitagorici che abbiamo citato in cui si afferma che la menzogna e l'invidia partecipano della natura dell'illimitato. Tuttavia una certa influenza deve avere pure avuto il fatto che esistevano paradossi legati all'infinito attuale che dimostravano le difficoltà logiche di tale concetto. Forse il paradosso più famoso è quello, di “Achille e la tartaruga” dovuto a Zenone.

Paradosso di Achille e la tartaruga. In questo paradosso si racconta di una sfida di una tartaruga (simbolo della lentezza) ad Achille (noto per la sua velocità) in una corsa. In tale sfida la tartaruga dichiara che purché gli siano dati dieci metri di vantaggio, non si sarebbe fatta raggiungere da Achille. Achille accetta la sfida, partono ed Achille percorre quei dieci metri di vantaggio. Tuttavia nel frattempo la tartaruga percorre un metro; Achille non si scoraggia ed allora percorre quel metro, ma nel frattempo la tartaruga percorre un decimetro; Achille percorre quel decimetro, ma nel frattempo la tartaruga percorre un centimetro . . . e così via all'infinito. In questo modo Achille può correre per sempre senza raggiungere mai la tartaruga.⁹

In definitiva, stante l'impossibilità per i Greci di estendere la nozione di numero naturale in quella di numero reale¹⁰, essi

⁹Attualmente tale paradosso viene "risolto" coinvolgendo la nozione di serie convergente. Infatti gli infiniti intervalli impiegati ogni volta da Achille per raggiungere la tartaruga diventano non solo sempre più piccoli ma il limite della loro somma *converge*. Di questa possibile soluzione (che comunque non è universalmente accettata ed anche a me non sembra centrare il problema) era convinto anche Cartesio come mostra il seguente passo in cui, riferendosi al paradosso afferma che:

"non è difficile a risolversi, quando si consideri che alla decima parte di una quantità viene aggiunta la decima di questa decima, e cioè una centesima; e poi ancora la decima di quest'ultima, ossia una millesima della prima; e così di seguito all'infinito, tutte queste decime prese insieme, benché siano supposte realmente infinite, non compongono tuttavia che una quantità finita. Ché se taluno dice che una tartaruga, la quale ha dieci leghe di precedenza rispetto a un cavallo dieci volte più veloce di lei, non potrà mai essere superata da questo, perché mentre il cavallo compie le dieci leghe la tartaruga ne percorre una e, mentre il cavallo supera questa lega, la tartaruga procede ancora di un decimo di lega e così all'infinito, bisogna rispondere che veramente il cavallo non la sopravvanzerà finché esso farà quella lega, quel decimo, quel centesimo, quel millesimo ecc. di lega; ma che non ne segue che non la supererà mai, perché quel decimo, centesimo, millesimo ecc. non fanno che un nono di lega, in capo al quale il cavallo comincerà a sopravvanzarla" (Lettres de M^r Descartes, Paris, 1657, N. 118).

¹⁰ D'altra parte in loro era totale la convinzione che gli unici numeri esistenti fossero i numeri naturali cioè gli interi positivi. Gli stessi numeri

furono indotti a considerare il continuo non riducibile alla nozione di numero e quindi a ritenere che :
la geometria è una scienza autonoma dalla aritmetica (in un certo senso la più importante tra le scienze).

5. Punti linee e Platonismo

Come abbiamo già detto, dalla scoperta delle grandezze incommensurabili in poi la geometria assume un ruolo centrale nella conoscenza scientifica e filosofica dell'antica Grecia. Essa ha uno sviluppo che appare enorme se lo si confronta con le altre branche della conoscenza. Per rendersene conto basta pensare che la geometria che si impara a scuola è solo una piccolissima parte della geometria scoperta di greci (se si esclude la geometria analitica che è una scoperta relativamente recente). Al contrario, gli argomenti di fisica, chimica, biologia che fanno parte dei programmi scolastici sono enormemente superiori per quantità e qualità a quelli che la persona più istruita dell'antica Grecia poteva possedere.

Esaminiamo gli aspetti più rilevanti della geometria dei Greci e partiamo da quello sicuramente più importante: l'idealizzazione degli enti geometrici.

Primo processo di idealizzazione. Un primo processo di idealizzazione consiste nel dare carattere di "sostanza" a quelle che prima erano considerate proprietà della materia. Se nella mate-

razionali erano considerati a volte come operatori, a volte come relazioni tra grandezze. Ad esempio quello che per noi è il numero $3/4$, per i greci non era un ente matematico in qualche modo esistente ma solo un modo abbreviato per dire "prendi un grandezza, moltiplicala per 3 e dividila per 4". Altre volte $3/4$ stava ad indicare una certa relazione tra due grandezze, infatti aveva senso scrivere, date due grandezze a e b , che a e b sono nella proporzione di 3 a 4, in breve $a:b = 3:4$ (si veda la teoria delle proporzioni). Un tale modo di vedere i razionali comportava poi difficoltà a definire le usuali operazioni di addizione e moltiplicazione. Infatti sembrava difficile giustificare l'addizione o la moltiplicazione di due operazioni o di due relazioni. A maggior ragione per i greci era inconcepibile una teoria degli irrazionali come quella attuale. Essi avevano però una tecnica, la teoria delle grandezze omogenee, che, come vedremo, permetteva ugualmente di esprimere quei concetti che, per i matematici moderni, coinvolgono gli irrazionali.

matica pre-ellenica "essere quadrato" era un attributo di alcuni oggetti materiali, non diverso da "essere rosso", "essere pesante", nel seguito si perverrà ad un nuovo ente "il quadrato" che a tutti gli effetti verrà trattato come una sostanza individuale. Questo significa che un quadrato diviene un oggetto di cui è possibile descrivere le proprietà allo stesso modo di come viene fatto per tutte le cose esistenti in natura. Dal punto di vista grammaticale, questo fenomeno si manifesterà nella trasformazione del ruolo di parole come "quadrato", "punto", "segmento" le quali da attributi divengono soggetti. Così a frasi del tipo "quel tavolo è un quadrato", che pongono in relazione un ente materiale (quel tavolo) con una sua possibile proprietà (essere quadrato) che appartengono, per così dire, alla fisica, si vengono sostituendo espressioni del tipo "il quadrato ha le diagonali uguali". Tali espressioni pongono in relazione enti e proprietà ideali e la loro validità, non potendo dipendere dalla esperienza del mondo esterno, può essere stabilita solo all'interno di una organizzazione razionale delle conoscenze.

Secondo processo di idealizzazione. Un secondo processo di idealizzazione è strettamente legato alla scoperta delle grandezze incommensurabili. Come abbiamo visto, ci si era accorti che un segmento non può essere costituito da una sequenza finita di punti materiali come pretendevano i pitagorici. D'altra parte se un segmento contiene quanti punti si vuole, allora tali punti devono necessariamente avere lunghezza nulla. Infatti, se per assurdo tutti i punti avessero grandezza l allora il segmento in questione dovrebbe avere lunghezza pari alla somma di infinite volte l e cioè dovrebbe avere lunghezza infinita. In conclusione:

i punti devono essere enti senza grandezza,¹¹

e, per analoghe considerazioni,

le linee devono essere enti senza larghezza

le superfici enti senza spessore.

La idealizzazione degli enti geometrici assume allora un aspetto radicale, in quanto nel mondo reale ogni cosa ha lunghezza, lar-

¹¹ Non è detto che i punti debbano essere assunti come concetti primitivi. In proposito si veda il mio articolo alla fine del capitolo.

ghezza e spessore. Naturalmente anche tutti gli altri enti geometrici come la sfera, il cubo, il cilindro sono astratti: non sarà mai possibile trovare in natura o costruire un corpo perfettamente sferico. Ma mentre il supporre l'esistenza di un corpo perfettamente sferico non sembra crearci grandi problemi, il supporre l'esistenza di qualcosa, il punto, che sia senza dimensioni è in completo contrasto con la concezione che abbiamo della materia. Siamo in presenza di un più alto livello di astrazione. In definitiva

- la concezione del punto come ente senza dimensioni appariva come logica conseguenza della scoperta degli incommensurabili
- i punti (più in generale, le linee, le superfici) non appartengono al mondo reale.

La conclusione a cui si doveva allora necessariamente pervenire era che:

si può avere conoscenza solo del mondo delle idee.

Il mondo percepito attraverso i sensi è qualcosa di illusorio al quale spetta solo il compito di "assomigliare" al mondo delle idee, così come un granello di sabbia può solo assomigliare ad un punto, non essere un punto. Un tale punto di vista assumerà la sua forma più compiuta nelle teorie filosofiche di Platone e sarà una delle cause dello scarsissimo sviluppo della fisica presso i Greci. Vediamo cosa dice Platone:

"I geometri si servono di figure visibili e ragionano su di esse, ma non ad esse pensando, bensì a ciò di cui quelle sono immagini, ragionando sul quadrato in sé e sulla diagonale in sé, e non su quella che disegnano. Lo stesso si dica per tutte le figure, che essi modellano e disegnano, di cui si servono come immagini (a guisa di ombre e di immagini riflesse nelle acque) cercando di vedere i veri enti che non si possono vedere se non col pensiero" (Platone).¹²

Ma come essere sicuri della validità di una asserzione geometrica dal momento che, non riferendosi al mondo reale, la sua verifica non può essere sperimentale? I greci osservarono che di alcune asserzioni si ha una intuizione talmente immediata che tutti sono

¹² Quindi è il mondo reale ad essere una immagine (sbiadita) del mondo delle idee e non il contrario, come saremmo portati a pensare ora. In altre parole, attualmente se una teoria non descrive bene la realtà, allora è la teoria che viene considerata non adeguata e non certo la realtà.

d'accordo sulla loro validità. Un esempio di tale tipo è *"per due punti distinti passa una sola retta"*. Vi sono però asserzioni, come il teorema di Pitagora, la cui validità non sembra essere altrettanto immediata. L'unica possibilità appariva allora quella di ricavare, tramite opportuni ragionamenti, le asserzioni più complesse da un gruppo prefissato di asserzioni sulla cui validità ci fosse un accordo generale: tali asserzioni venivano chiamate in alcuni casi postulati, in altri assiomi. Nasce in tale modo il metodo assiomatico che troverà negli elementi di Euclide la sua applicazione più bella e completa.

6. Gli elementi di Euclide.

Gli "Elementi" di Euclide rappresentano forse la tappa più importante dello sviluppo del pensiero scientifico moderno. L'opera, che consiste in 13 libri, fu composta da Euclide verso il 300 a.C. e rappresenta una organica esposizione di buona parte della matematica preesistente. Di questi libri i primi sei sono dedicati alla geometria piana, il settimo, l'ottavo, il nono ed il decimo sono di natura aritmetica, i rimanenti trattano di geometria solida.

Si inizia con una serie di definizioni, le prime quattro sono:

DEFINIZIONI

1. *Il punto è ciò che non ha parti*
2. *Una linea è una lunghezza senza larghezza*
3. *Estremi di una linea sono punti*
4. *Linea retta è quella che giace ugualmente rispetto ai suoi punti.*

Una critica che a volte è stata fatta a tali definizioni è che rimandano il concetto da definire ad altri concetti che non si sono definiti. Infatti per definire il punto si ricorre alla nozione di "parte", per definire le linee si utilizzano le nozioni di "lunghezza" e "larghezza", in 3 si usa la parola "estremi". Infine è alquanto oscuro che cosa significhi "giacere ugualmente". Invece attualmente una definizione viene fatta solo in funzione di nozioni già definite. Ad esempio data una struttura algebrica, un suo elemento e viene chiamato "elemento neutro" se risulta $x \cdot e = x$ e $e \cdot x = x$ per ogni elemento x . Tale definizione è ragionevole poiché utilizza solo la nozione di prodotto il quale, per il fatto che partiamo da una struttura algebrica data, risulta essere stata già definita. Tali criti-

che sono però ingiustificate perché per i greci le definizioni avevano un significato completamente diverso da quello attuale. Il loro ruolo (in un certo senso precedente ed esterno alla elaborazione scientifica) era quello di indicare, in qualche modo, enti che si riteneva avessero una esistenza propria e di cui ogni uomo ha una idea chiara. Prima di cominciare una trattazione scientifica di tali oggetti era infatti necessario, per potersi capire, essere sicuri che si stava parlando delle stesse cose. Ad esempio la Definizione 3 ci serve per capire che il termine "linea retta" significava quello che ora chiamiamo "segmento" e non ciò che oggi si intende per retta.¹³ In proposito possiamo parlare anche di "*definizioni reali*" nel senso che sono simili alle descrizioni-indicazioni che facciamo a volte di un oggetto esistente e conosciuto sia da noi che dalla persona con cui parliamo. Così "il punto è ciò che non ha parti" è una definizione allo stesso modo per cui lo è, ad esempio, "la penna di cui intendo parlare è quella posata sul tavolo". A volte si parla anche di "definizioni ostensive" per esprimere il fatto che esse servono solo ad indicare, a mostrare l'oggetto di cui si parla. Il punto di vista attuale, come vedremo quando esamineremo il metodo assiomatico, è totalmente diverso. Non si definisce ciascun ente isolatamente ma una intera classe di strutture ciascuna costituita da elementi matematici. Ad esempio non si propone una definizione di punto e di retta, piuttosto si definiscono delle strutture chiamate "spazi geometrici" i cui elementi base sono punti e rette. Oppure, le definizioni hanno il ruolo di indicare particolari elementi di una struttura.¹⁴

Negli Elementi di Euclide abbiamo poi una serie di *nozioni comuni* che riguardano l'uguaglianza o l'ordinamento.

¹³ D'altra parte, quando nella teoria degli insiemi un insegnante dice che un insieme è una "collezione", un "aggregato" di elementi, egli si comporta in modo analogo a quello di Euclide perché non spiega cosa si debba intendere per aggregato o collezione.

¹⁴ Ad esempio, in algebra, abbiamo sia la definizione di che cosa si debba intendere, ad esempio, per monoide, sia, dato un monoide, la definizione di che cosa sia un elemento neutro.

NOZIONI COMUNI

1. cose che sono uguali alla stessa cosa sono uguali tra loro
2. cose che coincidono tra loro sono uguali
3. se cose uguali sono addizionate a cose uguali, le totalità sono uguali.
4. se cose uguali sono addizionate a cose disuguali le totalità sono disuguali
5. il tutto è maggiore della parte.

La proprietà simmetrica per cui se $A = B$ allora $B = A$ viene data per scontata e non esplicitata. Il primo assioma dice che la uguaglianza è una relazione transitiva. Infatti afferma che da $A = C$ e $B = C$ (equivalentemente $C = B$) segue $A = B$. Il secondo assioma esprime la proprietà riflessiva. Pertanto i primi due assiomi, insieme alla proprietà simmetrica, ci dicono che $=$ è una relazione di equivalenza. Il terzo assioma dice che l'eguaglianza è quella che attualmente viene chiamata una congruenza, cioè una relazione di equivalenza compatibile con la struttura matematica che si considera (in questo caso la struttura additiva).¹⁵ L'ultimo assioma esprime la compatibilità dell'operazione $+$ con la relazione d'ordine.

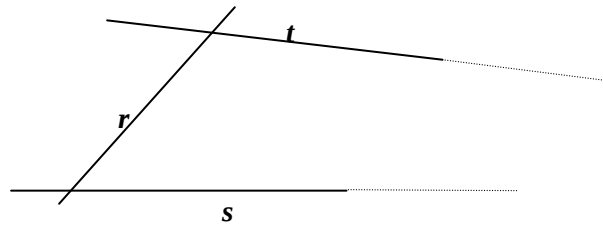
Infine seguono i cinque famosi postulati.

POSTULATI

- I** Risulti postulato che si possa condurre una linea retta da una qualsiasi punto ad ogni altro punto.
- II** e che una retta finita si possa prolungare continuamente in linea retta
- III** e che si possa descrivere un cerchio con qualsiasi centro ed ogni distanza
- IV** e che tutti gli angoli retti siano uguali tra loro
- V** e che, se una retta venendo a cadere su due rette forma due angoli interni e dalla stessa parte minori di due retti, le due rette prolungate illimitatamente verranno ad incontrarsi da quella parte in cui sono gli angoli minori di due retti.

¹⁵ Da notare che l'esigenza di affermare esplicitamente che cose che coincidono sono uguali mostra che l'uguaglianza non coincide in generale con l'identità, cioè che possono esistere cose uguali ma non identiche. D'altra parte due triangoli vengono definiti uguali se hanno lati uguali. Quindi due triangoli possono essere uguali senza coincidere (in quanto situati in parti diverse del piano).

Il quinto postulato è illustrato dalla seguente figura



ed afferma che se si prolungano i segmenti s e t prima o poi tali segmenti si incontreranno. Esso è noto come "assioma delle parallele" e, come vedremo, giocherà un ruolo importante nello sviluppo del pensiero matematico.

Le nozioni comuni si differenziano dai postulati per il fatto di non appartenere esclusivamente alla geometria ma a tutte le scienze. Inoltre, in un certo senso, esse hanno un grado di certezza maggiore dei postulati. Da notare che sia le nozioni comuni che i postulati saranno chiamati dai matematici posteriori con il nome di *assiomi*. I matematici greci erano molto più prudenti poiché né il termine "nozione comune" né quello di "postulato" hanno pretese universali. Con il primo termine si indicava qualche cosa accettata da tutti i membri di una comunità. Con il secondo termine solo qualche cosa di cui si "chiede" (e da ciò deriva il termine "postulato") una accettazione al proprio interlocutore per poter rendere possibile la successiva trattazione.

Evitare infinito ed illimitato. Come abbiamo già osservato, il rifiuto dell'infinito attuale porta i matematici greci ad evitare accuratamente il ricorso ad enti infiniti o illimitati, cioè all'infinito attuale. Così al posto di quelle che per noi sono le rette illimitate i greci si riferiscono costantemente ai segmenti. La illimitatezza (attuale) della retta si traduce, nel secondo postulato, nella indefinita prolungabilità dei segmenti, e lo stesso quinto postulato si riferisce a prolungamenti di segmenti. Non si deve poi pensare che i segmenti venissero considerati come oggetti infiniti (insiemi infiniti di punti). Infatti se è vero che dagli assiomi di Euclide è possibile dedurre che in ogni segmento giacciono "quanti punti si vuole", per i matematici greci questo non significava che un segmento è un insieme infinito di punti. Piuttosto segmenti e punti erano enti completamente indipendenti tra i quali sussisteva o meno una relazione di "giacenza" da non confondersi in

nessun modo con l'attuale relazione di appartenenza della teoria degli insiemi. Il fatto che, dato un segmento, se si trovano n punti giacenti in esso sia possibile trovarne anche $n+1$ non comporta l'accettazione dell'infinito attuale più di quanto lo comporti il fatto che, dato il numero intero n esista anche il numero $n+1$.

Carattere costruttivo dei postulati. Una ultima osservazione circa i postulati riguarda il loro carattere costruttivo. Si vede infatti chiaramente come essi siano corrispondenti ad operazioni che un disegnatore può eseguire, come il tracciare rette e cerchi. Fa eccezione forse il quinto postulato perché se la somma dei due angoli interni è minore di due retti solo per una quantità piccolissima, il verificare che le due rette in questione si incontrano comporta la necessità di tracciare segmenti più lunghi di quanto l'uomo è in grado di fare.

In definitiva potremmo anche dire che i postulati consentono di costruire, a partire da enti "a portata di mano" ancora enti "a portata di mano". Tuttavia il quinto assioma sembra comportarsi in modo diverso in quanto si riferisce al fatto che un punto di incontro tra le due rette prima o poi verrà trovato anche se (per dirla in modo un po' rozzo) questo punto di incontro fosse tanto lontano da non potere effettivamente essere costruito. Per questo motivo, o forse anche per altri, il quinto postulato venne accettato malvolentieri ed i matematici, da Euclide fino alla prima metà dell'ottocento, cercarono costantemente di eliminarlo dimostrandolo a partire dagli altri assiomi. Si deve comunque sottolineare che non veniva messo in discussione il fatto che tale postulato fosse vero. Solo che, per questioni di correttezza scientifica e di eleganza sarebbe stato opportuno ricorrere solo a postulati totalmente evidenti.

Il quinto postulato attualmente viene sostituito dal postulato:

Postulato delle parallele: data una retta r ed un punto P fuori da essa esiste al più una retta per P parallela ad r

che si prova essere equivalente. Si tenga presente che l'esistenza della retta parallela può essere dimostrata a partire dai rimanenti postulati.

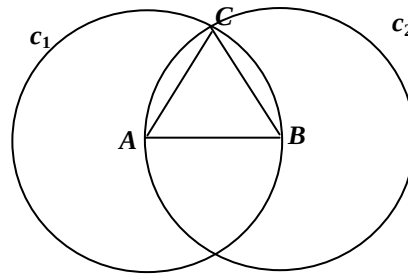
Concludiamo questo paragrafo con la prima delle proposizioni dimostrate da Euclide.

Proposizione 6.1. E' possibile, dato un segmento AB , costruire un triangolo equilatero con un lato uguale ad AB .

Dim Si tracci la circonferenza c_1 di centro A ed apertura AB e la circonferenza c_2 di centro B ed apertura BA e sia C il punto di incontro di tali circonferenze.

Allora il triangolo ABC è equilatero. Infatti AC è uguale ad AB in quanto entrambi raggi del cerchio c_1 e CB è uguale ad AB in quanto raggi del cerchio c_2 .

□



Come si vede si tratta di una dimostrazione semplice ed elegante ..., purtroppo proprio in questa che è la prima delle dimostrazioni proposte da Euclide si presenta un errore ! Infatti si utilizza, senza che sia stato dimostrato, il fatto che le due circonferenze si incontrano in un punto C . Per potere dimostrare questo fatto è necessario un postulato ulteriore, detto *postulato di continuità*, che negli Elementi viene a volte utilizzato anche se non viene messo nella lista dei postulati.¹⁶

¹⁶ Naturalmente si potrebbe anche pensare di trovare una dimostrazione diversa e corretta di tale proposizione, cioè una dimostrazione che non usi la proprietà di continuità. Purtroppo una tale dimostrazione non può esistere. Infatti consideriamo un modello di geometria in cui i punti sono gli elementi del prodotto cartesiano $Q \times Q$, cioè le coppie di numeri razionali ed in cui una retta è l'insieme dei punti che soddisfano una data equazione di primo grado a coefficienti razionali. Allora non è difficile controllare che in tale modello tutti i postulati di Euclide sono verificati (si veda anche il primo paragrafo del capitolo 3). Tale modello allora verificherà anche tutti teoremi che seguono da tali assiomi ed in particolare il teorema di Pitagora. Se allora si potesse dimostrare da tale sistema di assiomi la Proposizione 1, tale proposizione dovrebbe essere verificata dal nostro modello. Purtroppo ciò non è vero. Infatti consideriamo i punti $A = (-1,0)$ e $B = (1,0)$, e supponiamo che esista C tale che ABC sia equilatero. Allora è ovvio che tale punto dovrebbe essere del tipo $(0,q)$ e che, perché sia verificato il teorema di Pitagora, q^2 dovrebbe essere uguale a 3. Poiché non esiste nessun razionale il cui quadrato sia 3, ciò è assurdo.

7. La teoria delle grandezze omogenee (al posto dei numeri reali)

Le nozioni di classe di grandezze omogenee e di proporzione tra grandezze omogenee si trovano nel libro V degli *Elementi* di Euclide e sembrano risalire a Eudosso di Cnido, vissuto pochi decenni prima di Euclide. Tali nozioni permettono ai greci di fare molte delle cose che attualmente vengono fatte ricorrendo ai numeri reali positivi. Negli *Elementi* non viene proposto un sistema di assiomi completo per la nozione di classe di grandezze omogenee e spesso Euclide utilizza proprietà dettate dall'intuizione. Invece noi proviamo a proporre il seguente sistema di assiomi.

Definizione 7.1. Una *classe di grandezze omogenee* è una struttura $(G, =, <, +)$, tale che

A1 $+$ è una operazione commutativa ed associativa.

A2 per ogni $a \in G$ e per ogni $n \in \mathbb{N}$ esiste $u \in G$ tale che $n \cdot u = a$ (assioma della divisibilità).

A3 $<$ è una relazione d'ordine stretto totale compatibile con $+$.

Naturalmente si suppone che l'eguaglianza $=$ verifichi le proprietà elencate nelle nozioni comuni e che quindi sia una equivalenza compatibile con $+$ e $<$. Due grandezze che appartengono alla stessa classe di grandezze omogenee si dicono *omogenee* tra loro.

La classe dei numeri naturali soddisfa A1 e A3 ma non A2. Esempi di classe di grandezze omogenee sono i seguenti:

- la classe dei razionali positivi
- la classe dei reali positivi

Un esempio che più esprime l'idea dei greci di classe di grandezze omogenee è dato dall'insieme dei pesi di una bilancia. Chiamiamo uguali due pesi A e B se posti sui due piatti della bilancia rimangono in equilibrio. Diciamo invece che A è maggiore di B se il piatto su cui si poggia A si abbassa. La definizione di somma di due pesi è ovvia. Se vogliamo che l'assioma della divisibilità sia verificato dobbiamo accettare che dato un peso abbiamo a disposizione anche tutti i suoi sottomultipli.

Oltre agli assiomi elencati sono importanti anche il *Postulato di Archimede* e l'*Assioma della continuità* che presentano un interesse notevole anche nella matematica moderna ed in particolare nella teoria dei campi ordinati.

Postulato di Archimede. Siano u e b due grandezze omogenee con $u < b$, allora esiste un intero m tale che $m \cdot u > b$.

Tutti gli esempi numerici precedenti verificano l'assioma di Archimede. Per dare una idea di un comportamento non archimedeo, sia R l'insieme dei numeri reali e consideriamo la classe R^R delle funzioni di R in R . Inoltre definiamo la somma di due funzioni e l'ordinamento $\leq$ tra funzioni in modo usuale. In tale modello è subito visto che l'assioma di Archimede non è verificato. Infatti sia f una funzione limitata e g una funzione non limitata. Allora tutti i multipli di f sono ancora limitati e quindi non può esistere un multiplo di f maggiore di g . Il fatto che tale tipo di struttura non è archimedeo comporta che non risulta possibile misurare le funzioni utilizzando una funzione come unità di misura come invece è possibile fare per i segmenti.

Da notare che in questo esempio sono verificati tutti gli assiomi per le classi di grandezze omogenee tranne il fatto che l'ordinamento non è totale. Un esempio migliore di struttura non archimedeo è quello dei razionali non standard che considereremo nel seguito.

L'assioma della continuità afferma, in un certo senso, che non esistono "buchi" in una classe di grandezze omogenee.¹⁷ Esso comunque, pur essendo spesso utilizzato, non fu mai enunciato esplicitamente dai matematici greci. Bisogna aspettare Dedekind nel 1872, per avere una sua esplicita formulazione.

Definizione 7.2. Chiameremo *separati* due sottoinsiemi A e B di una classe di grandezze omogenee se ogni elemento di A è minore di ogni elemento di B . Un *elemento di separazione* è un elemento che è maggiorante di A e minorante di B .

Assioma di continuità (o di completezza). Ogni coppia A e B di insiemi separati ammette un elemento di separazione.

Si osservi che tale assioma non è verificato dalla classe $Q^+ = \{q \in Q : q > 0\}$ dei numeri razionali positivi. Infatti è possibile prova-

¹⁷ Abbiamo già accennato all'esigenza di un postulato di continuità nell'esaminare la dimostrazione del primo dei teoremi dimostrati negli *Elementi*.

re che le due classi $A = \{r \in \mathbb{Q}^+ \mid r^2 < 2\}$ e $B = \{r \in \mathbb{Q}^+ \mid r^2 > 2\}$ sono separate ma che non esiste nessun razionale q che sia elemento di separazione.

Classe di grandezze omogenee e processo di misurazione.

Il sistema di assiomi che abbiamo elencato ha lo scopo di rendere possibile una misurazione di una grandezza b rispetto ad una unità di misura u qualunque. Ad esempio assumiamo che u sia un regolo lungo un decimetro e che vogliamo misurare un segmento b . Allora la presenza di una operazione di somma consente di effettuare multipli $1 \cdot u, 2 \cdot u, \dots, m \cdot u$ di u cioè di riportare più volte consecutivamente u . Allo stesso momento, utilizzando il fatto che in una classe di grandezze omogenee è definita una uguaglianza ed un ordine stretto, possiamo man mano verificare se il multiplo $m \cdot u$ sia uguale o minore a b . Ad un certo momento è possibile che si ottenga un multiplo $m \cdot u$ di u coincidente proprio con b . In questo caso “fortunato” diremo che m è la misura di b rispetto all’unità di misura u . Se ciò non accade, allora per il Postulato di Archimede ad un certo punto si finisce comunque col superare b . Esiste cioè un numero m tale che $(m-1) \cdot u < b < m \cdot u$. Allora possiamo assumere i numeri $m-1$ ed m rispettivamente come misura per difetto e misura per eccesso di b rispetto a u . Se poi vogliamo ottenere una misura migliore, possiamo sostituire u con un suo sottomultiplo ad esempio porre $u' = u/10$ uguale al centimetro, cosa questa che l’assioma di divisibilità consente di fare. In questo caso, se esiste un intero m tale che $m \cdot u' = b$, cioè tale che $(m/10) \cdot u = b$ diciamo che la misura di b rispetto a u è il numero razionale $m/10$. Altrimenti, detto m un intero tale che $(m-1) \cdot u' < b < m \cdot u'$ cioè tale che $((m-1)/10) \cdot u < b < (m/10) \cdot u$ diremo che b ha come misura per difetto il razionale $(m-1)/10$ e misura per eccesso il razionale $m/10$.¹⁸ Se si vogliono misure più precise si procede poi dividendo il segmento u' in ulteriori 10 parti ottenendo un regolo di un millimetro e così via.

Non cambia il procedimento se invece che misure di segmenti si tratta di pesare oggetti. Se b è l’oggetto da pesare ed u (il grammo) un oggetto che funge da una unità di misura, allora si procede in questo modo. Si mette b in un piatto B ed ovviamente B si abbassa. Nel frattempo sull’altro piatto A si mettono man

¹⁸ Quest’ultimo passaggio naturalmente non poteva essere esplicitato dai greci per il fatto che essi non accettavano i razionali.

mano dei pesi uguali ad u . Se dopo avere messo m pesi uguali ad u la bilancia si equilibra, allora possiamo concludere che la misura di b rispetto ad u è esattamente m . In caso contrario ad un certo momento (assioma di Archimede) il piatto B si alza e possiamo concludere che il numero m dei pesi in A è una misura per eccesso di b , mentre $m-1$ ne è una misura per difetto. Utilizzando sottomultipli di u si possono ottenere delle pesate più precise.

Per fare invece un esempio del ruolo giocato dall'assioma della continuità, consideriamo il seguente problema:

data una circonferenza C trovare un segmento la cui lunghezza sia uguale a quella di C .

Allora posso considerare l'insieme A dei segmenti che si ottengono inscrivendo una poligonale in C e poi "raddrizzandola" in un segmento. In altre parole A è l'insieme dei segmenti la cui lunghezza è uguale alla lunghezza di una poligonale inscritta. Definiamo inoltre l'insieme B come l'insieme dei segmenti la cui lunghezza si ottiene "raddrizzando" una poligonale circoscritta. Precisamente B è l'insieme dei segmenti la cui lunghezza è uguale alla lunghezza di una poligonale circoscritta. Allora A e B sono due classi separate e pertanto, per l'assioma di continuità, esiste un segmento che separa tali classi. Tale segmento rappresenta il segmento la cui la lunghezza è uguale a quella della circonferenza.

8. La teoria delle proporzioni (al posto delle operazioni)

Abbiamo già parlato di incommensurabilità quando abbiamo provato che il lato e la diagonale del quadrato non sono commensurabili. In questo paragrafo riprendiamo tale nozione e la estendiamo nella teoria delle proporzioni.

Definizione 8.1. Due grandezze omogenee a e b si dicono *commensurabili* se esistono due interi n ed m tali che $n \cdot a = m \cdot b$.

In termini moderni noi diremmo che a e b sono commensurabili se $a = (m/n)b$, ma i Greci come abbiamo già osservato, non potevano fare riferimento ai razionali. Quello che è per noi il numero razionale m/n rappresentava per loro una relazione che doveva essere espressa, se si voleva mantenere il dovuto rigore, solo in termini di interi tramite l'uguaglianza $n \cdot a = m \cdot b$. Ad esempio Euclide nei suoi Elementi afferma che

“un rapporto è una sorta di relazione tra dimensioni di due grandezze della stessa specie.”

Proposizione 8.2. Due grandezze a e b sono commensurabili se e solo se hanno un sottomultiplo in comune cioè se esiste u tale $n \cdot u = a$ e $m \cdot u = b$ dove n ed m sono opportuni naturali.

Si passa poi alla definizione della nozione di proporzionalità tra quattro grandezze. Anche in questo caso si utilizza solo la nozione di multiplo e non quella di divisione. Ora se attualmente volessimo esprimere il fatto che quattro grandezze a, b, c, d (che in generale possono essere non razionali) sono in proporzione, cioè che il numero reale $a : b$ è uguale al numero reale $c : d$, potremmo farlo ricorrendo solo ai razionali dicendo che

- un razionale n/m è minore di $a:b$ se e solo se è minore di $c:d$
- un razionale n/m è maggiore di $a:b$ se e solo se è maggiore di $c:d$.

D'altra parte possiamo riscrivere tali equivalenze facendo uso solo dei numeri interi dicendo che:

- $nb \leq ma$ se e solo se $nd \leq mc$
- $nb \geq ma$ se e solo se $nd \geq mc$.

Ciò suggerisce la seguente definizione (due o più grandezze vengono chiamate omogenee tra loro se appartengono alla stessa classe di grandezze omogenee).

Definizione 8.3. Siano a, b, c, d quattro grandezze con a omogeneo a b e c omogeneo a d . Diremo che tali grandezze sono in proporzione e scriveremo $a:b = c:d$ se per ogni coppia di interi n, m risulta che

$$nb \leq ma \Leftrightarrow nd \leq mc \quad \text{e} \quad nb \geq ma \Leftrightarrow nd \geq mc.$$

Si noti che non è necessario che le quattro grandezze siano omogenee tra loro, è sufficiente assumere che a sia omogenea a b e c a d ; cioè che a e b appartengano ad una stessa classe di grandezze omogenee e c e d ad un'altra classe di grandezze omogenee. Ad esempio è possibile parlare di proporzione anche nel caso in cui a e b sono lunghezze e c e d aree.

Da osservare che l'uso dell'uguaglianza nell'espressione $a:b = c:d$ non deve far ritenere che i Greci ritenessero che $a:b$ e $c:d$ fossero "oggetti" uguali. Infatti a e b non sono da considerare numeri reali di cui si effettua la divisione e quindi $a:b$ da solo non denota

niente. La equazione $a:b = c:d$ era considerata come un modo per indicare una relazione tra quattro grandezze e non una uguaglianza tra due. Il segno di uguaglianza può verificarsi soltanto nel caso di coppie di grandezze commensurabili che sono le uniche che ammettono multipli comuni. Tuttavia sembra che Euclide consideri una scrittura del genere anche come una sorta di affermazione per cui la coppia (a,b) ha qualcosa in comune con la coppia (c,d) (si veda la nota successiva).

Proposizione 8.4. Per l'uguaglianza dei rapporti valgono le proprietà riflessiva, simmetrica e transitiva (cioè è una relazione di equivalenza tra coppie).¹⁹

Enunciamo ora il teorema di esistenza del quarto proporzionale.

Teorema 8.5. (Teorema di esistenza del quarto proporzionale). Siano a, b due grandezze appartenenti alla classe M di grandezze omogenee, e c una grandezza appartenente alla classe N , allora esiste una grandezza $x \in N$ tale che $a:b = c:x$.

Il teorema di esistenza del quarto proporzionale equivale all'affermazione che il prodotto e la divisione sono definite nella classe delle grandezze omogenee. Infatti, considerando il caso $a = 1$, $1:b = c:x$ equivale a dire che (con il linguaggio attuale della teoria dei numeri reali) $x = b \cdot c$. Se si pone $b = 1$ allora $a:1 = c:x$ equivale a dire che $c = ax$.

9. Misure, equiscomponibilità, equicompletabilità

¹⁹ Se si assume il punto di vista della matematica moderna, il fatto che valgano le proprietà riflessiva, simmetrica e transitiva induce a considerare una nuova classe di oggetti che si possono costruire al modo seguente. Data una classe G di grandezze omogenee definiamo in $G \times G$ la relazione $\equiv$ ponendo $(a,b) \equiv (c,d)$ se a, b, c, d sono in proporzione. Poiché una tale relazione è di equivalenza, ripartisce $G \times G$ in classi di equivalenza. Chiamiamo *rapporto* una classe completa di equivalenza. In questo modo è lecito dire che le due classi a e b hanno lo stesso rapporto per dire che appartengono alla stessa classe di equivalenza. Da questo punto di vista il simbolo di uguaglianza nell'espressione $a:b = c:d$ riacquista il suo significato usuale di identità.

La teoria delle proporzioni fornisce agli antichi greci uno strumento per elaborare metodi che corrispondono a quella che per noi è la “misurazione” di una figura geometrica (calcolo dell’area, volume, perimetro). Infatti, invece di procedere alla misurazione intesa come assegnazione di un numero reale ad una figura, essi si concentravano sul confronto tra le grandezze di figure geometriche. Questo confronto consisteva nel constatare una uguaglianza o una proporzione. Ad esempio, invece di dire che l’area del triangolo è uguale alla base per l’altezza diviso due, veniva affermato che l’area di un triangolo sta all’area di un rettangolo con la stessa base e la stessa altezza come 1 sta a 2. Per quello che riguarda l’uguaglianza, una nozione che spesso veniva utilizzata era quella di “equiscomponibilità”. Detto in termini intuitivi, tale nozione può essere definita al modo seguente.

Definizione 9.1. Due figure geometriche F e $\underline{F}$ si dicono *equiscomponibili* se è possibile

Tagliare F nelle figure $X_1, \dots, X_n$,

Spostare tali figure in modo da ottenere le figure $\underline{X}_1, \dots, \underline{X}_n$

Ricomporre $\underline{X}_1, \dots, \underline{X}_n$ in modo da ottenere $\underline{F}$.

In altre parole due figure sono equiscomponibili se sono la “somma” di figure uguali. Naturalmente è necessario precisare che cosa si intenda per “figura geometrica”, “tagliare”, “spostare”, “ricomporre”. Accettiamo per ora che per figura geometrica si intenda un qualunque insieme di punti. La nozione di “tagliare” può essere rappresentata dalla nozione insiemistica di partizione.

Definizione 9.2. Se X è un insieme chiamiamo *partizione finita* di X , una classe finita $X_1, \dots, X_n$ di insiemi tali che

- $X = X_1 \cup \dots \cup X_n$
- per ogni i e j con $i \neq j$, $X_i \cap X_j = \emptyset$.

Per la nozione “spostare” si può ricorrere alla nozione di isometria.

Definizione 9.3. Chiamiamo *isometria* del piano euclideo E ogni funzione $i : R^2 \rightarrow R^2$ che conserva le distanze, cioè tale che

$$d(x, y) = d(i(x), i(y)).$$

Diciamo che i *sposta* un insieme di punti A nell'insieme di punti B se $i(A) = B$. Due figure geometriche A e B si dicono *isometriche* o *uguali* se esiste una isometria che sposta A in B .

Esempi tipici di isometrie nel piano sono le traslazioni, le rotazioni, ed i ribaltamenti rispetto ad un asse.

Possiamo ora dare la nozione di equiscomponibilità in modo più preciso.

Definizione 9.4. Due figure geometriche F e $\underline{F}$ si dicono *equiscomponibili* se ammettono rispettivamente due partizioni $X_1, \dots, X_n$, e $\underline{X}_1, \dots, \underline{X}_n$, con X_i isometrico a $\underline{X}_i$.

La seguente proposizione mostra l'importanza della equiscomponibilità per il calcolo delle aree.

Teorema 9.5. Se due figure F e $\underline{F}$ sono equiscomponibili allora hanno la stessa misura.²⁰

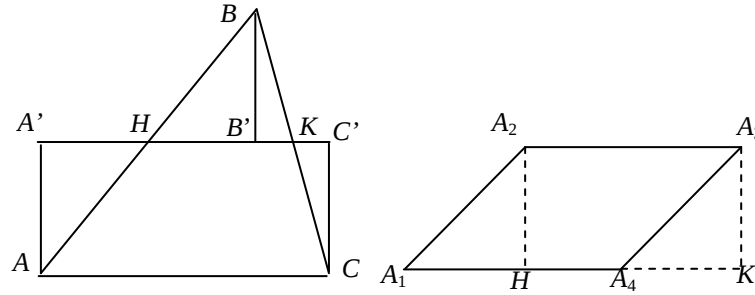
Dim. Per la finita additività della misura, la misura di F è la somma delle misure di $X_1, \dots, X_n$, e la misura di $\underline{F}$ è la somma delle misure di $\underline{X}_1, \dots, \underline{X}_n$. D'altra parte se si accetta che le misure rimangono invariate durante gli spostamenti, allora le misure di $X_1, \dots, X_n$, coincidono con le misure di $\underline{X}_1, \dots, \underline{X}_n$. $\square$

Due famose applicazioni di tale teorema corrispondono al calcolo dell'area di un triangolo ed a quella di un parallelogramma.

²⁰ Tale proposizione presuppone che la scomposizione delle due figure avvenga tramite pezzi che siano "misurabili", cioè dotati di area. Senza tale ipotesi la proposizione non vale (si veda il teorema 10.1 per un controesempio). Il fatto che tutte le figure geometriche fossero misurabili era dato per scontato dai matematici greci. D'altra parte le figure prese in considerazione erano solo triangoli, rettangoli cerchi, parallelogrammi o simili, figure che si possono ottenere componendo tali figure elementari. Figure di questo tipo sono tutte misurabili.

Teorema 9.6. Ogni triangolo è equiscomponibile ad un rettangolo con la stessa base ed avente come altezza la metà dell'altezza del triangolo. Ogni parallelogramma è equiscomponibile ad un rettangolo che ha la stessa base e la stessa altezza.

H



Dim. Dato un triangolo di vertici A, B, C (si veda la figura a sinistra), tagliamo a metà i lati AB e BC nei punti H e K , tracciamo il rettangolo $A'C'AC$, e l'altezza BB' del triangolo HBK . L'equiscomponibilità tra il triangolo ABC ed il rettangolo $A'C'AC$ segue dal fatto che tali figure hanno il quadrilatero $AHKC$ in comune, il triangolo HBB' è uguale a $AA'H$ ed il triangolo $B'BK$ è uguale a $KC'C$. Possiamo visualizzare tale equiscomponibilità immaginando di tagliare HBK nei due triangolini e di fare ruotare tali triangolini intorno ai punti H e K .

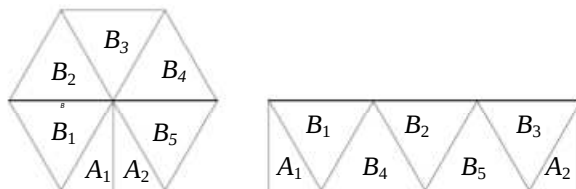
Per la seconda parte della proposizione riferiamoci alla figura disegnata a destra. E' facile dimostrare che il triangolo A_1A_2H è uguale al triangolo A_4A_3K . Pertanto il parallelogramma $A_1A_2A_3A_4$ si può scomporre nel triangolo A_1A_2H ed il quadrilatero $HA_2A_3A_4$ i quali, opportunamente ricomposti costituiranno il rettangolo A_2HKA_3 avente la stessa base e la stessa altezza. $\square$

Attualmente questi risultati si esprimono dicendo che l'area del triangolo è uguale alla base per l'altezza diviso due e l'area del parallelogramma è uguale alla base per l'altezza.

Problema. Dimostrare, utilizzando la nozione di equiscomponibilità, che l'area di un trapezio è data dalla somma delle basi per l'altezza ed il prodotto diviso due.

Teorema 9.7. Un poligono regolare è equiscomponibile ad un rettangolo che ha come base la metà del perimetro e come altezza l'apotema.

Dim. Riferendoci ad esempio ad un esagono, è sufficiente considerare la seguente figura:



in cui si mostra che l'esagono si può scomporre in 7 triangoli che ricomposti formano un rettangolo. $\square$

Attualmente questo teorema si esprime dicendo che l'area di un poligono regolare è uguale al perimetro per l'apotema diviso due.

Abbiamo visto che in un certo senso due figure sono equiscomponibili se sono “somma” di pezzi uguali. Possiamo anche considerare la possibilità che due figure siano la “differenza” tra figure uguali.

Definizione 9.8. Due figure si dicono *equicompletabili* se aggiungendo ad esse parti a due a due uguali si ottengono poligoni uguali o equiscomponibili.

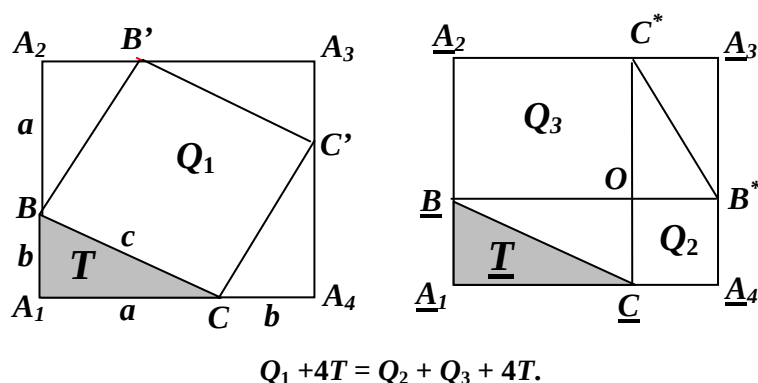
Proposizione 9.9. Due figure equicompletabili hanno la stessa misura.

Per mostrare un esempio di applicazione della nozione di equicompletabilità, esponiamo una dimostrazione del teorema di Pitagora differente da quella già esposta.²¹

²¹ La dimostrazione precedente è di carattere algebrico ed è più semplice da un punto di vista logico. Questa che ora esponiamo è di carattere geometrico ed ha il pregio di fare “vedere in un colpo d'occhio” la validità del teorema di Pitagora. Infatti riusciamo a percepire che Q_1 e $Q_2 + Q_3$ sono uguali guardando queste figure come differenze di figure uguali. Dimostrazioni di tale tipo in didattica della matematica prendono nome di dimostrazioni “visuali”.

Teorema 9.10. (Teorema di Pitagora) Dato un triangolo rettangolo, se si considera l'unione dei due quadrati costruiti sui cateti otteniamo una figura che è equicompletabile con il quadrato costruito sull'ipotenusa.

Dim. Ripetiamo la costruzione geometrica di un quadrato $Q = A_1A_2A_3A_4$ effettuata nella dimostrazione già fatta del teorema di Pitagora. Costruiamo un quadrato $\underline{Q} = \underline{A_1A_2A_3A_4}$ uguale a $Q = A_1A_2A_3A_4$ ed un triangolo $\underline{T} = \underline{A_1B\underline{C}}$ uguale al triangolo T . Tracciamo da $\underline{B}$ e $\underline{C}$ due perpendicolari ai lati del quadrato ed indichiamo con B^* e C^* i punti di intersezione con $\underline{A_4A_3}$ e con $\underline{A_2A_3}$, rispettivamente. Si perviene alle due figure seguenti:



Si ottengono due quadrati Q_2 e Q_3 (di lati uguali a quelli dei due cateti) e quattro triangoli uguali a T . Allora dall'uguaglianza di Q con Q^* , segue che l'unione dei due quadrati Q_2 e Q_3 è equicompletabile con il quadrato Q_1 . $\square$

10. L' equiscomponibilità è un metodo universale

In questo paragrafo vogliamo mostrare che l'equiscomponibilità è un metodo universale per il calcolo delle aree, cioè che se due figure piane a contorni rettilinei hanno la stessa area allora sono equiscomponibili. Per fare questo per prima cosa dimostriamo che l'equiscomponibilità è una relazione di equivalenza. Tale proprietà segue sostanzialmente dal fatto che l'insieme delle isometrie costituisce un gruppo.

Proposizione 10.1. L' equiscomponibilità è una relazione di equivalenza, cioè è riflessiva simmetrica e transitiva.

Poniamoci ora il problema se la proposizione 9.2 si possa invertire, cioè se tutte le volte che due figure hanno la stessa area allora sono equiscomponibili. La seguente bella proposizione, dimostrata da Bolyai nel 1832 e dal dilettante di matematica Gerwin (1833), fornisce una risposta positiva alla nostra domanda.

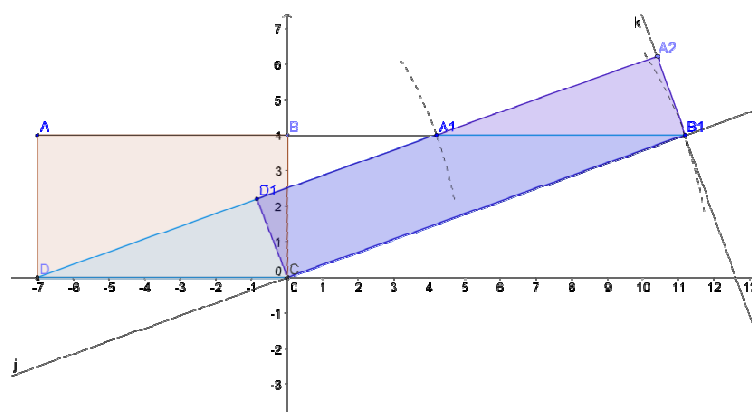
Teorema 10.2. Due figure poligonali che sono di uguale misura sono equiscomponibili.²²

Dim. Cominciamo con l'osservare che dalla seconda parte della proposizione 9.6 si ricava che:

1. *due parallelogrammi che hanno una base uguale e stessa area sono equiscomponibili.*

Infatti tali parallelogrammi hanno necessariamente anche la stessa altezza e quindi, come dimostrato nel teorema 9.4, sono entrambi equiscomponibili ad uno stesso rettangolo. Per la proprietà transitiva sono quindi equiscomponibili tra loro.

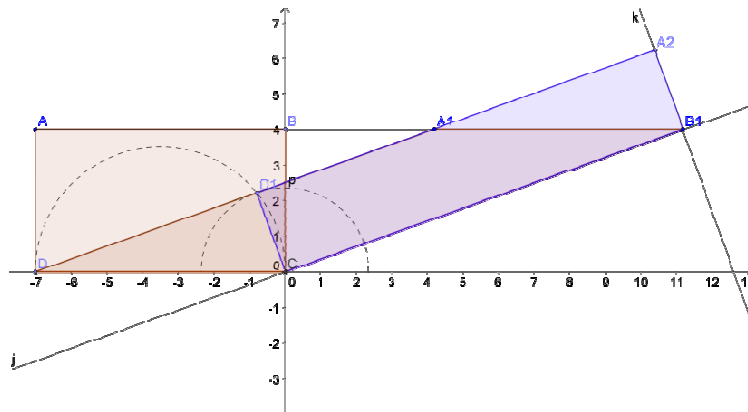
2. *Dato un rettangolo ABCD, comunque si fissi un segmento EF esiste un rettangolo di base EF equiscomponibile ad ABCD.*



²² Da tale teorema segue che due poligoni equicompletabili, avendo uguale misura, sono anche equiscomponibili. Poiché è evidente che due poligoni equiscomponibili sono equicompletabili, risulta che le due nozioni coincidono.

Supponiamo EF maggiore di CB , allora, tenendo ferma la base DC “deformiamo” il rettangolo $ABCD$ in un parallelogramma A_1B_1CD con A_1 e B_1 appartenenti alla retta AD ed in modo che il lato CB_1 sia uguale ad EF .²³ Successivamente “rad-
drizziamo” tale parallelogramma tenendo fisso il lato CB_1 e lasciando immutata l’altezza rispetto CB_1 . Si ottiene un rettangolo $CB_1D_1A_2$. Allora, essendo $ABCD$ equiscomponibile al parallelogramma A_1B_1DC ed essendo tale parallelogramma equiscomponibile con il rettangolo $CB_1D_1A_2$, $ABCD$ risulta equiscomponibile con $CB_1D_1A_2$.

Supponiamo EF minore di CB ed indichiamo con D_1 l’intersezione del cerchio di diametro DC con il cerchio di centro C e raggio EF . L’angolo in D_1 è retto. Chiamiamo A_1 l’intersezione di BL con AD ed B_1 l’intersezione di BL con la parallela per C alla retta DD_1 . Come nel caso precedente avremo che il rettangolo $ABCD$ è equiscomponibile con il parallelogramma A_1B_1DC che a sua volta è equiscomponibile al rettangolo $CB_1D_1A_2$.

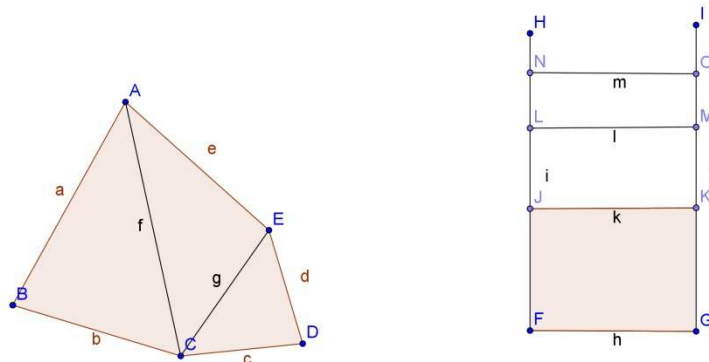


Proviamo infine che:

3. Data una figura poligonale Q ed un segmento FG esiste un rettangolo di base FG equiscomponibile a Q .

²³ Graficamente, è sufficiente intersecare la retta AB con i cerchi di raggio EF e centro D e C ed ottenendo i punti A_1 e B_1 rispettivamente. Il rettangolo $CB_1D_1A_2$ si ottiene in modo ovvio.

Infatti la poligonale può essere ripartita in un numero finito di triangoli $T_1, \dots, T_n$ e tali triangoli sono equiscomponibili a rettangoli $R_1, \dots, R_n$ che possiamo supporre tutti aventi una base uguale ad FG . E' sufficiente allora incollare tali rettangoli lungo tali basi per ottenere un rettangolo equiscomponibile a Q (si veda la figura successiva in cui Q è un pentagono che si divide in tre triangoli).



Concludiamo osservando che se due poligoni Q_1 e Q_2 hanno misura uguale allora, fissato un segmento AB esiste un rettangolo R_1 di base uguale ad AB equiscomponibile a Q_1 ed un rettangolo R_2 di base AB equiscomponibile a Q_2 . Poiché Q_1 e Q_2 hanno per ipotesi la stessa area, R_1 ed R_2 hanno la stessa area e quindi sono uguali. In conclusione essendo Q_1 e Q_2 equiscomponibili a rettangoli uguali sono equiscomponibili tra loro. $\square$

E' interessante osservare che, data una figura polinomiale Q , se il rettangolo equiscomponibile a Q lo si costruisce con un lato di lunghezza unitaria, allora la lunghezza dell'altro lato rappresenta proprio l'ampiezza dell'area. Un'altra osservazione è che le costruzioni effettuate nella dimostrazione del teorema sono tutte effettuabili con riga e compasso. Ciò nel senso che l'espressione "esiste una partizione" può essere intesa nel senso che i punti dove "tagliare" una figura sono tutti costruibili con riga e compasso. Ad esempio nel dire che un triangolo è equiscomponibile ad un rettangolo, il rettangolo viene costruito considerando i punti medi di due lati. D'altra parte i punti medi di un segmento possono, appunto, essere trovati con il compasso.

Terzo problema di Hilbert Si pone la questione se un teorema simile non possa essere dimostrato anche nella geometria dello spazio. Tale problema fu incluso al terzo posto della famosa lista di problemi proposti da Hilbert nel 1900. La risposta è negativa: nel 1903 il matematico Dehn mostrò che esistono due tetraedri di uguale volume che non sono equiscomponibili.

11. Contro i matematici

Quando ho detto che la matematica giocava un ruolo centrale nella cultura dei greci non intendevo dire che ciò accadeva per tutti i greci. In realtà, essendo gli antichi greci un popolo notevolmente intelligente, vivace e non conformista, accadeva che su di un dato argomento ciascuno avesse una propria idea. Così sulla matematica non tutte le opinioni erano concordi e lo stesso Aristotele non era convinto della importanza della matematica allo stesso modo di Platone.

Uno di quelli che meno avevano in simpatia la matematica era Sesto Empirico, un filosofo vissuto attorno al 180 dopo Cristo. Nel suo libro "Contro i geometri" egli critica in maniera radicale l'opera dei matematici. Ad esempio critica le basi logiche del metodo ipotetico deduttivo di Euclide e precisamente la validità di un tale metodo come strumento per ottenere verità sul mondo.

Infatti per Sesto Empirico il valore conoscitivo di un sistema ipotetico-deduttivo non può stare nelle ipotesi (cioè nei postulati) in quanto o una asserzione è vera, ed allora non ha senso considerarla come ipotesi, oppure è falsa ed allora è sbagliato prenderla come ipotesi.

“ . . . la cosa ammessa per ipotesi o è vera e tale come noi la supponiamo, o è falsa. Ma se essa è vera, noi non la postuliamo, perché non sentiamo il bisogno di ricorrere ad una cosa piena di sospetto quale è l'ipotesi, ma l'assumiamo immediatamente, giacché nessuno assume ipoteticamente le cose vere ed esistenti, quali ad esempio il fatto che adesso è giorno ed io sto discutendo e respirando . . . Se però essa non è tale, ma è falsa, non si ricava alcun vantaggio dell'ipotesi . . . ”

Si potrebbe allora dire che il valore di un sistema ipotetico deduttivo stia nella validità delle conseguenze che se ne ricavano.

veri. Un caso è dato dalla teoria degli insiemi che ha prodotto molti ed utili teoremi nonostante che la scoperta dei paradossi abbia dimostrato la sua falsità. D'altra parte ogni teoria fisica è stata dimostrata essere falsa dalla teoria successiva, in un certo senso.

Altre obiezioni sono inerenti direttamente alla matematica e precisamente alla concezione degli enti ideali.

" . . . essi dicono, inoltre, che una linea viene prodotta dallo scorrimento di un punto, una superficie dallo scorrimento di una linea, e un corpo solido dallo scorrimento di una superficie . . . "

" . . . il punto che essi definiscono come segno-privo-di-dimensioni, si deve concepire o come corporeo o come incorporeo. Corpo esso non è, secondo le loro stesse affermazioni, giacché le cose che non hanno dimensione, secondo loro, non sono corpi. Resta allora da dire che esso è incorporeo, il che è ancora una volta incredibile. Infatti ciò che è incorporeo non si può concepire come generatore di una linea; quindi il punto non è un segno-privo-di-dimensioni. "

In altri termini Sesto si pone il problema di come un ente senza dimensioni, il punto, possa generare (per scorrimento) un ente con una dimensione, la linea. Tale osservazione equivale, in un certo senso, ad osservare che se la lunghezza di un punto è zero allora ogni segmento, in quanto insieme (somma) di punti, deve avere lunghezza zero. La differenza consiste nel riferirsi alla linea prodotta dallo scorrimento di un punto (operazione questa che non sembra coinvolgere l'infinito attuale) e non alla linea intesa come insieme di punti (concezione questa che, coinvolgendo l'infinito attuale, non era presa in considerazione). Tale tipo di argomentazione viene ripetuta anche per confutare il concetto di linea senza larghezza il cui scorrimento genera una superficie e quello di superficie senza spessore il cui scorrimento genera i solidi.

Infine un altro tipo di critica riguarda il procedimento di astrazione mediante il quale l'uomo perverrebbe a concepire gli enti matematici ideali. Infatti per Sesto tutto ciò che viene concepito viene concepito o mediante una diretta esperienza oppure tramite un procedimento di immaginazione-astrazione. Ora è evidente che non possiamo mai avere una esperienza diretta, ad

esempio, di una linea senza larghezza e che questa idea di linea non è simile a niente di esistente

. . . giacché non cade sotto i nostri sensi una lunghezza che sia priva di larghezza . . .

D'altra parte un procedimento di immaginazione-astrazione può avvenire

. . . per somiglianza, ad esempio dall'immagine di Socrate lo stesso Socrate, per composizione, ad esempio dal cavallo e dall'uomo un ippocentauro giacché mescolando le membra del cavallo e dell'uomo noi siamo riusciti ad immaginare l'ippocentauro che non è né uomo né cavallo ma è composto da entrambi. Per analogia, infine, si concepisce qualcosa ancora in due guise ossia o per accrescimento o per diminuzione, come quando, ad esempio, tenendo presenti gli uomini normali . . . concepiamo per accrescimento il Ciclope . . . e come quando per diminuzione immaginiamo un pigmeo che non ci è mai caduto sotto i sensi.

Ora è evidente che il concetto di linea senza larghezza non è simile a niente di esistente per lo stesso motivo per cui non è determinato dall'esperienza. E' anche immediato che tale concetto non si ottiene per composizione come nel caso dell' ippocentauro. Non resta altro che il procedimento di diminuzione ma anche questo permette solo di ridurre a piacere (potenzialmente) la larghezza non di considerarla (attualmente) nulla e quindi non permette il tipo di astrazione che si richiederebbe. Anche in questo caso entra in gioco il rifiuto dell'infinito attuale. Ad esempio se noi volessimo definire un punto P immaginando una piccola sfera s_1 di centro P , poi un'altra sfera s_2 di centro P e raggio dimezzato e così via, allora il punto P sarebbe il frutto del processo di astrazione definito in tale modo solo se noi potessimo considerare tale processo terminato con un'operazione di limite che sarebbe possibile solo se si accettasse l'infinito attuale.

Un altro tipo ancora di procedimento possibile di astrazione è quello che fa pervenire ad un concetto mediante una semplice cancellazione di alcune delle proprietà dell'oggetto di partenza. Così il concetto di linea senza larghezza si può ottenere semplicemente "facendo finta" che la larghezza di un oggetto reale non esiste. Ma:

. . . se noi, dopo aver concepito una certa lunghezza avente una data quantità di larghezza, abbiamo altresì la possibilità di assumere una lunghezza priva di larghezza sopprimendo quest'ultima, allora allo stesso modo, dopo aver concepito un pezzo di carne che abbia la proprietà di essere vulnerabile, mediante la soppressione di tale proprietà noi potremo anche concepire una carne che non sia soggetta alla vulnerabilità . . . Ma tale cosa è completamente impossibile e contraria alle comuni nozioni umane: infatti ciò che viene concepito come invulnerabile secondo noi non è affatto carne, giacché la carne, in quanto carne, viene concepita con la proprietà di essere vulnerabile . . . Onde anche la lunghezza concepita come priva di larghezza non potrebbe essere una lunghezza, giacché la lunghezza, in quanto lunghezza, viene concepita come avente una certa quantità di larghezza.

In altre parole se un oggetto ideale $\hat{A}$ si ottiene da un oggetto concreto A facendo astrazione da alcune particolari proprietà allora non è chiaro perché si possa considerare $\hat{A}$ un rappresentante di A (o viceversa), cioè non è chiaro che relazione sussiste tra A ed $\hat{A}$. In particolare, supponiamo di avere dimostrato una proposizione per l'ente ideale $\hat{A}$, allora chi ci autorizza a dire che tale proposizione è valida anche per A ? Ad esempio supponiamo che un astronomo debba studiare l'orbita di un asteroide A in presenza dei pianeti $A_1, A_2, \dots, A_n$ e che per fare questo decida di rappresentare con il punto $\hat{A}$ l'asteroide e con i punti $A_1', \dots, A_n'$ i pianeti (su cui si concentra la massa). Immaginiamo inoltre che, applicando la geometria euclidea e la meccanica del moto dei punti, giunga ad una qualche conclusione P . Chi ci assicura che P sia valida oltre che per $\hat{A}$ anche per l'asteroide reale A ?

LETTURA

Platone e la duplicazione del quadrato (da *IL MENONE*).²⁴

MENONE: Va bene, caro Socrate, ma in che senso tu sostieni che noi non apprendiamo ma che ciò che noi chiamiamo “apprendimento” è reminiscenza? Sapresti insegnarmi che è veramente così?

SOCRATE: Già prima dicevo, caro Menone, che sei un furbacchione, ora mi domandi se so insegnarti proprio mentre sto dicendo che non c'è insegnamento ma reminiscenza evidentemente per farmi subito apparire in contraddizione con me stesso.

MENONE: No, per Zeus, o Socrate, non l'ho detto con questo scopo, ma solo per l'abitudine. Se, però, in qualche modo mi puoi dimostrare che la cosa sta così come dici, allora dimostromelo.

SOCRATE: Non è facile! Tuttavia, per te, sono disposto a farlo. Chiamami un po' uno dei tuoi numerosi servi che sono qui, quello che vuoi tu, perché tramite lui ti possa dare la dimostrazione.

MENONE: Certo. Vieni qui, ragazzo!

SOCRATE: E' greco e parla greco?

MENONE: Sì, perfettamente. E' nato in casa.

SOCRATE: Fa' bene attenzione, se ti sembra che si ricordi o che impari da me.

MENONE: Presterò attenzione.

SOCRATE: Dimmi un po', ragazzo, sai che questa qui è un'area quadrata (abcd)?

²⁴ In questo dialogo si dimostra un metodo per duplicare un quadrato. La dimostrazione equivale, in un certo senso, alla dimostrazione del teorema di Pitagora nel caso particolare di un triangolo rettangolo i cui cateti sono uguali. Infatti se i cateti misurano c e l'ipotenusa misura i , allora il fatto che $2 \cdot c^2 = i^2$ equivale a dire che il quadrato costruito sull'ipotenusa è il doppio del quadrato costruito su un cateto. Da notare che la duplicazione del cubo è un problema notevolmente più complesso ed è stato dimostrato che non può essere affrontato con i soli strumenti della riga e compasso. Tale dialogo, in cui si vuole giustificare la teoria della reminiscenza, costituisce un bellissimo esempio di “didattica della matematica” in quanto lo scopo di Socrate è quello di fare “emergere” la dimostrazione del teorema dalla mente del servo.

RAGAZZO: Sì.

SOCRATE: Il quadrato è dunque una superficie che ha uguali tutti questi lati, che sono quattro (ab, bc, cd, da).



RAGAZZO: Certamente.

SOCRATE: E non ha forse uguali anche queste linee qui, che lo attraversano nel mezzo (ac, bd)?

RAGAZZO: Sì.

SOCRATE: E non potrebbe esserci forse una superficie come questa e più grande e più piccola?

RAGAZZO: Certamente.

SOCRATE: Se dunque questo lato (bc) fosse di due piedi, e anche questo (ab) di due, di quanti piedi sarebbe l'intero? Fa questa considerazione: se da questa parte (ab) fosse di due piedi e da quest'altra (bc) di uno solo, la superficie non sarebbe forse di una volta due piedi?

RAGAZZO: Sì .

SOCRATE: Ma, poiché anche da questa parte (bc) è di due piedi, non diventa di due volte due piedi?

RAGAZZO: Sì, diventa.

SOCRATE: Diventa, perciò, di due volte due piedi?

RAGAZZO: Esatto .

SOCRATE: E quanti sono, allora, due volte due piedi? Fa' il conto e dillo.

RAGAZZO: Quattro, o Socrate.

SOCRATE: E non potrebbe darsi un'altra superficie doppia di questa, ma tale da avere tutti i lati eguali come questa?

RAGAZZO: Sì.

SOCRATE: Di quanti piedi sarà dunque?

RAGAZZO: Di otto.

SOCRATE: E ora cerca di dirmi di quanto sarà ciascun lato di essa. Il lato di questa è di due piedi; e, allora, di quanto sarà quello di quella doppia?

RAGAZZO: E' chiaro, o Socrate, che sarà doppio.

SOCRATE: Vedi, caro Menone, che io non gli insegno, ma che lo interrogo su ogni cosa? Ed ora, costui ritiene di sapere quale sia il lato dal quale deriverà l'area di otto piedi: o non ti sembra?

MENONE: A me sì.

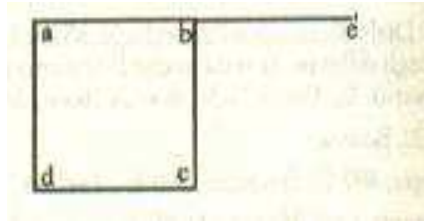
SOCRATE: E lo sa, dunque?

MENONE: Per nulla.

SOCRATE: Però ritiene che derivi dal lato doppio.

MENONE: Sì.

SOCRATE: Osserva come verrà via via ricordandosi, come appunto deve ricordarsi. E tu dimmi: dal lato doppio, dici che ha origine la superficie doppia? E tale, dico, che non sia di qui lunga e di qui corta, ma che sia eguale da ogni parte come questa qui, però doppia di questa, ossia di otto piedi. Ma sta' attento, se ti sembra ancora che possa derivare dal lato doppio.

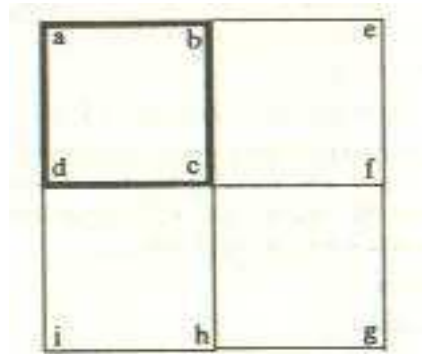


RAGAZZO: A me sì.

SOCRATE: E non diventa forse questo lato (ae) doppio di questo (ab), se ne aggiungiamo un altro come questo, da questa parte (be)?

RAGAZZO: Certamente.

SOCRATE: Da questo (ae), dici tu, deriverà la superficie di otto piedi, quando si tracceranno quattro lati come questi



RAGAZZO: Esattamente.

SOCRATE: Ma in questa superficie non ci sono forse queste quattro qui (abcd, befc, cfgh, dchi), delle quali ognuna è uguale a questa di quattro piedi (abcd)?

RAGAZZO: Sì.

SOCRATE: Disegniamo, allora, a partire da questo 4

lati uguali. E' oppure no questa la superficie (aegi) che tu affermi essere di 8 piedi?

RAGAZZO: Esattamente.

SOCRATE : Ma in questa superficie non vi sono forse queste 4 qui (abcd , befc , cgh , dchi), delle quali ognuna é uguale a questa di 4 piedi?

RAGAZZO : Sì

SOCRATE: E quanto diventa allora? Non diventa quattro volte questa?

RAGAZZO: E come no?

SOCRATE: E allora, è il doppio quattro volte tanto?

RAGAZZO: No, per Zeus.

SOCRATE: Ma quante volte?

RAGAZZO: Quadruplo.

SOCRATE: Dunque, dal lato doppio, o ragazzo, non deriva una superficie doppia ma quadrupla.

RAGAZZO: Dici il vero.

SOCRATE: E quattro volte quattro, fanno sedici, no?

RAGAZZO: Sì.

SOCRATE: E allora, quella di otto piedi da quale lato? Non se ne ottiene da questo (ae) una quadrupla?

RAGAZZO: Sì, lo dico.

SOCRATE: E quella di quattro, dalla metà di questo qui (ae)?

RAGAZZO: Sì.

SOCRATE: Ebbene, l'area di otto piedi non è forse doppia di questa qui (abcd), e metà di quest'altra (aegi)?

RAGAZZO: Sì.

SOCRATE: E allora, non deriverà da un lato maggiore rispetto a questo (ab), ma minore rispetto a quest'altro (ae); o no?

RAGAZZO: Così mi pare.

SOCRATE: Bene; quello che a te sembra devi rispondere. E dimmi: questo lato (ab) non era di due piedi e quest'altro (ae) di quattro?

RAGAZZO: Sì.

SOCRATE: Bisogna allora che il lato della superficie di otto piedi sia maggiore di questo di due, ma minore di quello di quattro.

RAGAZZO: Bisogna .

SOCRATE: Cerca allora di dire di che lunghezza tu affermi che esso debba essere.

RAGAZZO: Di tre piedi.

SOCRATE: Se dev'essere di tre piedi, aggiungiamo dunque a questo lato (ab) la metà di questo (ah), e avremo i tre piedi (ah). Questi sono due piedi (ah) e questo uno (hh). Alla stessa manie-

ra, a partire di qua si ottengono due piedi (ab) più un piede (dc).
Ne deriva, così, l'area che tu dici (abil).

RAGAZZO: Sì.

SOCRATE: Ma se da questa parte (ab) è di tre, e da quest'altra (hi) di tre, l'intera superficie non diventa di tre volte tre piedi?

RAGAZZO: Sembra.

SOCRATE: E tre volte tre, quante volte sono?

RAGAZZO: Nove.

SOCRATE: E il doppio, di quanti piedi doveva essere?

RAGAZZO: Otto.

SOCRATE: Dal lato di tre piedi non deriva per nulla la superficie di otto.

RAGAZZO: Certamente no.

SOCRATE: Ma allora da quale lato? Cerca di dircelo con esattezza; e se non vuoi fare calcoli, indicaci almeno da quale.

RAGAZZO: Ma per Zeus, o Socrate, io non lo so.

SOCRATE: Comprendi ora, o Menone, a che punto si trova attualmente nel processo del ricordare? Prima, cioè, non sapeva quale fosse il lato del quadrato di otto piedi, come del resto neppure ora lo sa; tuttavia, allora credeva di saperlo, e rispondeva con sicurezza come se sapesse e non riteneva di aver dubbi; ora è convinto di aver dubbi e come non sa, così neppure crede di sapere.

MENONE: Dici il vero.

SOCRATE: Non si trova dunque, ora, in una situazione migliore, relativamente alla cosa che non sapeva?

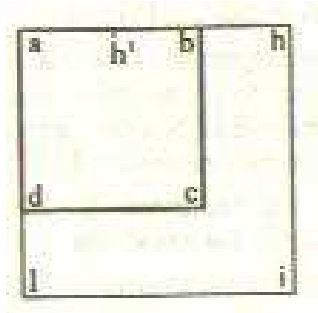
MENONE: Anche questo mi pare:

SOCRATE: Avendolo fatto dubitare, pertanto, e avendolo fatto intorpidire come fa la torpedine, gli abbiamo forse nuociuto?

MENONE: Non mi pare.

SOCRATE: Dunque, come sembra, gli abbiamo recato giovamento, al fine della ricerca di come stia effettivamente la cosa. Ora, infatti, ricercerebbe anche di buon grado, dal momento che non sa; mentre allora, facilmente, di fronte a molti e spesso avrebbe creduto di dire bene, affermando che per ottenere una superficie doppia, bisogna prendere il lato doppio in lunghezza.

MENONE: Sembra.



SOCRATE: Credi, dunque, che egli si sarebbe messo a cercare o ad imparare ciò che egli riteneva di sapere non sapendolo, prima che fosse caduto nel dubbio ritenendo di non sapere, e che avesse desiderato di conoscere?

MENONE: Non mi pare, o Socrate.

SOCRATE: Dunque, l'intorpidimento gli ha giovato?

MENONE: Mi sembra.

SOCRATE: Osserva, ora, da questo dubbio come scoprirà la verità, ricercando insieme a me, mentre io non farò altro che interrogarlo, senza insegnargli. . E fa bene attenzione che tu non mi colga ad insegnargli o a spiegargli, e non solo ad interrogarlo intorno alle sue convinzioni.

Dimmi, dunque: non è di quattro piedi questa superficie (abcd)? Comprendi?

RAGAZZO: Sì.

SOCRATE: Potremmo aggiungere ad essa quest'altra eguale (befc)?

RAGAZZO: Sì.

SOCRATE: E quest'altra terza, uguale a ciascuna di queste (cfgh)?

RAGAZZO: Sì.

SOCRATE: E non potremmo anche completare la figura in questo angolo (dchi)?

RAGAZZO: Certamente.

SOCRATE: E non risulteranno queste quattro superfici eguali ?

RAGAZZO: Sì.

SOCRATE: E, allora, tutto questo intero (aegi), quante volte diventa più grande di questo (abcd)

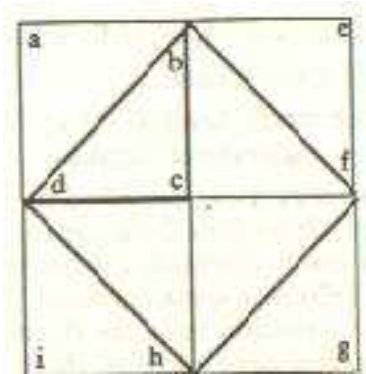
RAGAZZO: Quattro volte.

SOCRATE: Per noi, invece, doveva essere il doppio; o non ricordi?

RAGAZZO: Certamente.

SOCRATE: E questa linea tracciata da un angolo all'altro (bd, bf, fh, hd), non viene forse a dividere a metà ciascuna di queste superfici?

RAGAZZO: Sì.



SOCRATE: Non si ottengono, dunque, queste quattro linee uguali racchiudenti quest'area qui (bfhd)?

RAGAZZO: Sì, si ottengono.

SOCRATE: Considera allora: quanto grande è questa superficie (bfhd)?

RAGAZZO: Non lo so.

SOCRATE: Di questi quadrati, che sono quattro, ciascuna linea non ha tagliato internamente la metà di ciascuno? O no?

RAGAZZO: Sì.

SOCRATE: E quante ve ne sono di queste metà in questa figura (bfhd)?

RAGAZZO: Quattro.

SOCRATE: E quante in quest'altra (abcd)?

RAGAZZO: Due.

SOCRATE: E il quattro che cos'è rispetto al due?

RAGAZZO: Il doppio.

SOCRATE: Questa superficie, dunque, di quanti piedi diventa?

RAGAZZO: Di otto piedi.

SOCRATE: Da quale linea?

RAGAZZO: Da questa (a'b).

SOCRATE: Da quella che abbiamo tracciata da un angolo all'altro del quadrato di otto piedi?

RAGAZZO: Sì.

SOCRATE: Coloro che se ne intendono chiamano questa linea diagonale; sicché, se essa ha nome diagonale, allora dalla diagonale, come tu dici, o ragazzo di Menone, si può ottenere l'area doppia.

RAGAZZO: Certamente, o Socrate.

SOCRATE: Che cosa ti sembra, o Menone? C'è qualche pensiero da lui espresso che non sia suo?

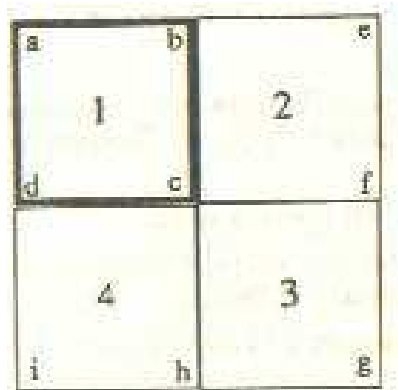
MENONE: No, tutti suoi.

SOCRATE: Eppure, non sapeva, come dicevamo poco fa.

MENONE: Dici il vero.

SOCRATE: E c'erano in lui questi pensieri o no?

MENONE: Sì.



SOCRATE: Dunque, in chi non sa intorno alle cose che non sa, vi sono opinioni vere che ad esse si riferiscono?

MENONE: Sembra.

SOCRATE: Ora in lui, come un sogno, sono state suscitate queste opinioni; e, interrogandolo di nuovo più volte e in molti modi su queste stesse cose, sta certo che finirà per sapere con precisione, sulle medesime, non meno esattamente di ogni altro .

MENONE: Pare proprio di sì.

SOCRATE: Dunque, egli saprà senza che nessuno gli insegni, ma solo che lo interroghi, traendo egli stesso la scienza da se medesimo.

RAGAZZO: Sì.

SOCRATE: E questo trarre la scienza di dentro a sé, non è ricordare?

MENONE: Certamente.

SOCRATE: E la scienza che ora egli possiede, o la imparò un tempo o la possedette sempre.

MENONE: Sì.

SOCRATE: Dunque, se la possedette sempre, fu anche sempre conoscente; e se, invece, l'ha appresa in un tempo, non poté certo averla appresa nella presente vita. Oppure gli insegnò qualcuno geometria? Costui, infatti, farà lo stesso per tutta la geometria, e per tutte quante le altre scienze. C'è, forse, uno che gli abbia insegnato tutto? A buon diritto tu devi saperlo: non per altro, perché è nato ed è stato allevato in casa tua.

MENONE: Ma lo so che nessuno gli ha mai fornito insegnamenti.

SOCRATE: Ed ha o non ha queste conoscenze?

MENONE: Necessariamente, o Socrate, sembra.

SOCRATE: Allora, se non le ha acquisite nella presente vita, questo non è ormai evidente, ossia che le ebbe e le apprese in un altro tempo?

MENONE: E' chiaro.

SOCRATE: E non è forse questo il tempo in cui egli non era uomo?

MENONE: Sì .

SOCRATE: Se, allora, e nel tempo in cui è uomo e nel tempo in cui non lo è, vi sono in lui opinioni vere, le quali, risvegliate mediante l'interrogazione, diventano conoscenze, l'anima di lui non sarà stata in possesso del sapere sempre in ogni tempo? E' evi-

dente infatti che, nel corso di tutto quanto il tempo, talora è e talora non è uomo.

MENONE: E' chiaro.

SOCRATE: Se, dunque, sempre la verità degli esseri è nella nostra anima, l'anima dovrà essere immortale. Sicché bisogna mettersi con fiducia a ricercare ed a ricordare ciò che attualmente non si sa: questo è infatti ciò che non si ricorda.

MENONE: Mi sembra che tu dica bene, o Socrate, ma non so come.

CAPITOLO 2

CRISI DELLA GEOMETRIA EUCLIDEA

Per tutto il diciassettesimo ed il diciottesimo secolo la geometria rimase, nella guerra contro l'empirismo, una fortezza inespugnabile degli idealisti. Coloro i quali credevano (come in generale si credeva sul continente) che fosse possibile una conoscenza del mondo reale certa ed indipendente dall'esperienza non avevano che da indicare la geometria: soltanto un pazzo avrebbe messo in dubbio la sua validità, e soltanto uno sciocco ne avrebbe negato il riferimento oggettivo. (Bertrand Russell)

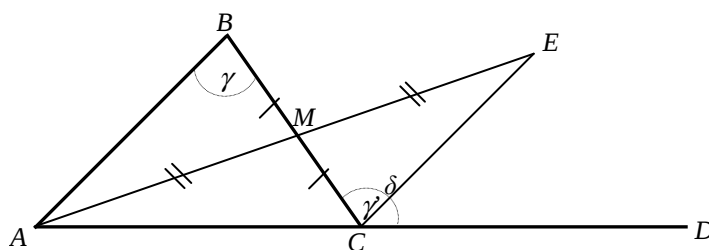
1. Crisi del carattere assoluto della geometria

La geometria euclidea restò al centro della matematica fino agli inizi del 1800 quando, ad opera di matematici come Lobachevsky e Bolyai si svilupparono le geometrie non euclidee, cioè geometrie che negavano validità al quinto postulato. Prima di parlare delle geometrie non euclidee, esaminiamo più da vicino il significato del quinto postulato. Per prima cosa mostriamo che dagli assiomi di Euclide escluso il quinto è possibile comunque derivare che data una retta r ed un punto P fuori da essa esiste almeno una retta parallela ad r passante per P . Il quinto postulato serve invece per dimostrare che tale retta è unica. Pur non volendo esporre le dimostrazioni complete, vediamo i passi fondamentali per provare l'esistenza della parallela e cominciamo con il Teorema dell'angolo esterno.

Proposizione 1.1. Un angolo esterno di un triangolo è sempre strettamente maggiore dei due angoli interni non adiacenti.

Dim. Consideriamo un triangolo ABC e proviamo ad esempio che l'angolo esterno $\delta = \angle BCD$ è maggiore dell'angolo $\gamma = \angle ABC$. A tale scopo denotiamo con M il punto medio del segmento BC e prolunghiamo il segmento AM in un segmento AE in modo che AM sia uguale a ME . Allora i due triangoli ABM e MEC sono uguali avendo i due angoli in M uguali in quanto opposti al vertice e due lati uguali per costruzione. In particolare γ sarà uguale all'angolo

$\gamma' = BCE$. D'altra parte tale angolo, essendo una parte di δ , risulta minore di δ . In conclusione γ è minore di δ .

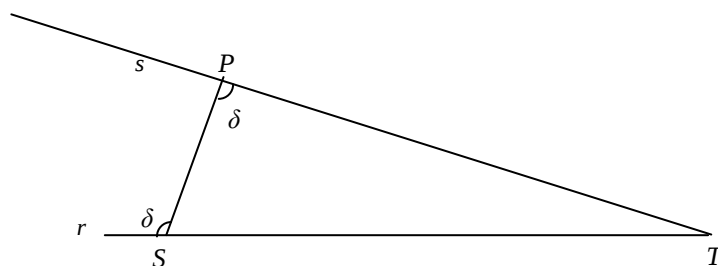


□

Naturalmente la dimostrazione che abbiamo dato diventa rigorosa solo se prima si è provato che tutte le costruzioni fatte sono rese possibili dai postulati di Euclide (escluso il quinto). Ad esempio deve essere prima provato che esiste il punto medio di un segmento, che due angoli opposti al vertice sono uguali e che valgono i criteri di uguaglianza dei triangoli. Uno discorso analogo vale per tutte le dimostrazioni di carattere euclideo che sono fatte in questo capitolo.

Proposizione 1.2. Data una retta r ed un punto P non appartenente ad r esiste almeno una retta per P parallela ad r .

Dim. Sia S un punto qualsiasi della retta r , tracciamo la retta per



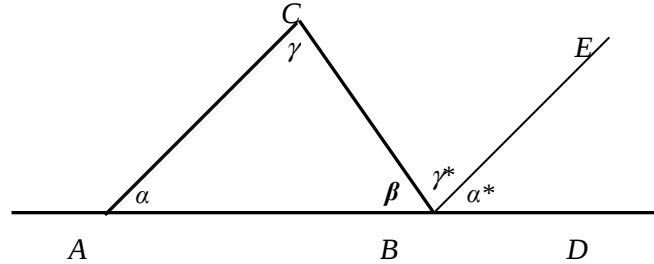
T

S e per P ed indichiamo con δ uno degli angoli tra r ed il segmento SP . Tracciamo poi una retta s per P che formi con SP l'angolo δ . Allora se per assurdo le rette r ed s si incontrassero in un punto T , nel triangolo SPT avremmo l'angolo esterno in S uguale all'angolo interno non adiacente in P , in contrasto con quanto dimostrato nella proposizione precedente. Quindi la retta s non incontra r . □

Se si accetta il quinto postulato allora è possibile dimostrare un teorema più generale del teorema dell'angolo esterno e questo anche in maniera più semplice.

Teorema 1.3. Un angolo esterno è uguale alla somma degli angoli interni non adiacenti. Ne segue che la somma degli angoli interni di un triangolo è un angolo piatto.

Dim. Dato il triangolo di vertici A, B, C tracciamo da B la



parallela BE alla retta AC . In tale modo si ottiene una coppia di rette parallele tagliate dalla trasversale CB . Per un teorema (che supponiamo di avere già dimostrato) sugli angoli alterni interni abbiamo che $\gamma = \gamma^*$ mentre per il teorema sugli angoli corrispondenti abbiamo che $\alpha = \alpha^*$. Ciò prova che l'angolo esterno CBD è uguale ad $\alpha + \gamma$. Ne segue anche che $\beta + \alpha + \gamma = \beta + \gamma^* + \alpha^*$ è un angolo piatto. $\square$

Da tale teorema ovviamente segue anche il teorema dell'angolo esterno. Tuttavia Euclide preferisce effettuare la dimostrazione del teorema dell'angolo esterno che abbiamo esposto in Proposizione 1.1 probabilmente perché non utilizza il postulato delle parallele. D'altra parte l'uso di tale postulato viene rimandato negli Elementi quanto più possibile quasi a dimostrazione di qualche perplessità. Naturalmente i matematici greci erano convinti della verità del quinto postulato. Il problema è che pare fossero anche convinti che esso fosse in realtà dimostrabile utilizzando gli altri assiomi della geometria. Tale convincimento rimase comune anche a tutti i matematici successivi fino al 1800. Infatti numerosi furono i tentativi di trovare una sua dimostrazione.

Tuttavia nella prima metà del 1800 accade un fatto nuovo: tre differenti matematici svilupparono un nuovo tipo di geometria. Si trattava del tedesco Gauss, dell'ungherese Bolyai e del russo

Lobachevsky che fecero la loro scoperta uno indipendentemente dall'altro. Punto di partenza di tali autori è una nuova visione della geometria che da essi viene considerata un ramo della fisica piuttosto che un prodotto a priori del nostro spirito. Ciò comportava che la validità degli assiomi della geometria dovesse essere convalidata o confutata dall'esperienza. Ad esempio ecco cosa scrive Gauss.

Secondo la mia più profonda convinzione, la teoria dello spazio ha nei confronti del nostro sapere una posizione completamente diversa da quella della pura teoria delle grandezze (aritmetica); infatti, viene assolutamente a mancare alla nostra conoscenza della prima quella completa convinzione della sua necessità (e quindi anche della sua assoluta verità), che invece inerisce alla seconda; dobbiamo umilmente ammettere che, mentre il numero è puramente un prodotto del nostro spirito, lo spazio possiede una realtà anche al di fuori del nostro spirito, alla quale noi non possiamo prescrivere le sue leggi completamente a priori. (Lettera di Gauss a Bessel del 9 aprile 1830)

Naturalmente il considerare la geometria alla stessa stregua della fisica comportava una interpretazione delle nozioni geometriche in termini di oggetti esistenti nel mondo reale. Ad esempio il segmento congiungente due punti P e Q poteva essere assimilato al percorso di minima distanza tra P e Q , oppure alla linea percorsa da un raggio di luce che partito da P raggiunge Q . Ciò permetteva ad esempio di immaginare esperimenti capaci di stabilire se un dato triangolo avesse o meno come somma di angoli interni un angolo piatto. A questo nuovo punto di vista nei confronti della geometria corrispondeva poi anche un atteggiamento diverso nei confronti del quinto postulato. Infatti se per i matematici greci il volere dimostrare tale postulato era solo una questione di eleganza e semplicità e nessuno ne metteva in dubbio la validità, i matematici dell'ottocento accettavano anche la possibilità che esso non fosse verificato nello spazio reale. Ecco quello che dice Lobacevskij nell'introduzione al suo libro "Nuovi principi della geometria" del 1835.

A tutti è noto che, fino ad oggi, nella geometria la teoria delle parallele è rimasta incompiuta. Gli sforzi inutili compiuti dai

tempi di Euclide, per il corso di duemila anni, mi spinsero a sospettare che nei suoi stessi concetti non si racchiude ancora quella verità che si voleva dimostrare, e che può essere controllata, in modo simile alle altre leggi fisiche, soltanto dall'esperienza quali, ad esempio, le osservazioni astronomiche.

Ma nel caso che il postulato delle parallele fosse falso nello spazio reale, allora esso non poteva essere certo provato a partire dai rimanenti postulati (certamente veri). Infatti da cose vere è possibile dedurre solo cose vere.

2. Modelli di geometrie non euclidee

Lo strumento per provare in maniera inconfutabile che il quinto postulato non è derivabile dai rimanenti è quello di esibire "modelli" matematici che non verificano il quinto postulato ma che verificano tutti i postulati rimanenti. Tali modelli furono chiamati *modelli non euclidei della geometria*. Esponiamo, sommariamente e solo per darne una idea, due modelli di geometria non euclidea, quello di Klein e quello di Poincaré.

Il modello di Klein. In tale modello

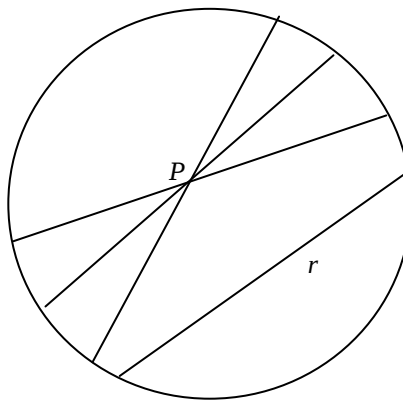
- i punti sono i punti interni ad un dato cerchio S di centro C ,
- le rette sono le corde del cerchio (esclusi gli estremi).

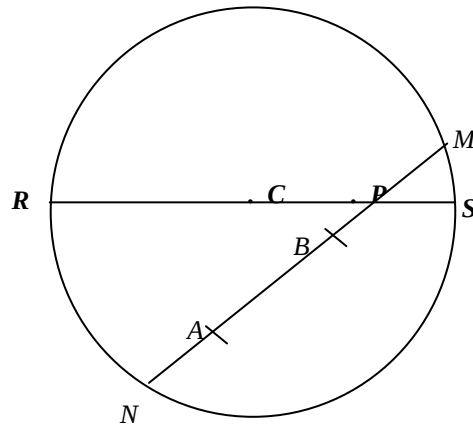
Come è illustrato nella figura, per un punto P fuori di una retta r passano infinite rette parallele ad r . Pertanto non vale l'assioma delle parallele. È immediato provare che tutti gli altri assiomi di Euclide sono verificati. Ad esempio, è evidente che per due punti passa una ed una sola retta.

Inoltre la nozione di distanza di due punti A e B (che serve per dare la nozione di eguaglianza tra segmenti) si definisce ponendo

$$d(A,B) = -\log\left(\frac{MB \cdot NA}{MA \cdot NB}\right)$$

dove M ed N sono i punti di intersezione di AB con la circonferenza.





Per capire il significato di tale distanza, supponiamo ad esempio che la circonferenza abbia equazione $x^2+y^2=1$, e calcoliamoci la distanza di un punto P di coordinate $(x,0)$ dal centro C .

$$d(C,P) = -\log\left(\frac{PS \cdot 1}{1 \cdot RP}\right) = -\log\left(\frac{(1-x)}{(1+x)}\right).$$

Questa formula è interessante perché mostra che:

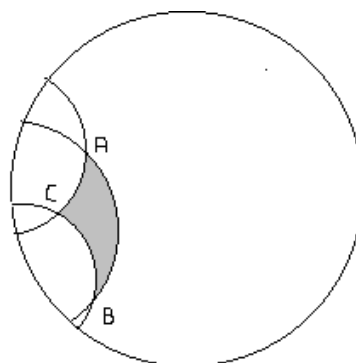
la distanza $d(C,P)$ tende a diventare infinita quando x tende ad 1, cioè quando P si avvicina al bordo del cerchio.

Ciò significa che un segmento unitario che venga spostato verso il bordo deve subire una contrazione che è tanto più forte quanto più ci si allontani dal centro. Equivalentemente possiamo dire che se, partendo dal centro del cerchio, cominciamo ad allontanarci a passi regolari, tali passi divengono sempre più piccoli senza che noi ce ne accorgiamo. Pertanto per un essere vivente all'interno del cerchio non risulta possibile uscire dal cerchio che gli apparirà, a tutti gli effetti, un universo non limitato. Ciò in contrasto con quanto appare ad un osservatore "esterno" per il quale il modello di Klein sembra occupare una parte limitata dello spazio. Possiamo immaginare una tale situazione supponendo che i punti perimetrali del cerchio esercitino una forza di repulsione verso i punti interni e che tale forza risulti tanto più forte quanto più ci si avvicina a tali bordi.

Si osservi che la nozione di uguaglianza di angoli viene definita in modo analogo a quanto fatto per l'uguaglianza di segmenti e che tale nozione non coincide con quella usuale del piano eucli-

deo. Inoltre, nonostante le apparenze, un triangolo di tale geometria ha somma degli angoli interni minore di un angolo piatto.

Modello di Poincaré. Consideriamo ancora come insieme di punti l'insieme dei punti interni ad una circonferenza S ma cambiamo la nozione di retta. Infatti chiamiamo “rette” tutte i diametri e tutte le circonferenze perpendicolari ad S . La nozione di lunghezza di un segmento si definisce in maniera non troppo diversa da quella del modello di Klein mentre la nozione di uguaglianza di angoli è quella usuale del piano euclideo. Ciò è interessante in quanto, come si vede nel triangolo ABC nella figura, permette di comprendere perché la somma degli angoli interni di un triangolo è minore di un angolo retto.



Sia il modello di Klein che quello di Poincaré permettono di dimostrare il seguente teorema:

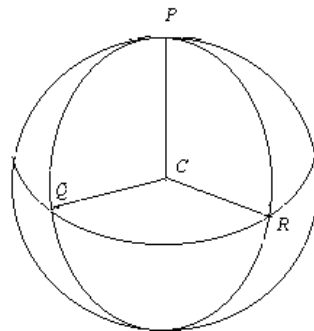
Teorema 2.1. Esiste un modello che verifica tutti gli assiomi di Euclide tranne il quinto postulato. Pertanto il quinto postulato non può essere dimostrato a partire dai rimanenti assiomi.

Dim. Indichiamo con T la teoria costituita da tutti gli assiomi proposti da Euclide escluso il quinto postulato. Allora se tale postulato fosse dimostrabile a partire da T esso sarebbe vero in tutti i modelli di T . Ma questo è impossibile poiché abbiamo trovato dei modelli di T che non verificano tale postulato.¹ $\square$

¹ Per capire lo schema di tale ragionamento consideriamo il caso semplice in cui T è la teoria dei gruppi e sia α la proprietà commutativa. Allora α non può essere un teorema di T perché in tale caso tutti i gruppi sarebbero commutativi. In altre parole il mostrare un esempio di gruppo che non è commutativo mostra l'indipendenza di α da T .

3. Altre geometrie.

Successivamente il matematico Riemann ed il fisico Helmholtz, uno indipendentemente dall'altro, svilupparono geometrie in cui la somma degli angoli interni di un triangolo è strettamente maggiore di un angolo piatto. In esse non veniva negato solo il quinto



postulato ma anche che un segmento possa essere prolungato a piacere. Per avere una idea di tali tipi di geometrie si osservi che nel piano euclideo i segmenti possono essere definiti come le linee di minima lunghezza congiungenti due punti assegnati. Ciò suggerisce di identificare le linee rette come lo strumento per potersi muovere con minore fatica possibile su di una superficie, e questo qualsiasi sia la superficie. Al variare delle superfici si ottengono geometrie diverse. Supponiamo ad esempio che ci si muova lungo la superficie di una sfera, allora appare ancora naturale chiamare "segmento" la linea più breve sopra tale superficie che congiunga due punti dati. È questo il punto di vista di un capitano di una nave che si muova sulla superficie terrestre e che, ovviamente, deve scegliere la rotta più breve per raggiungere la sua meta. Ora, se tale sfera ha centro C , si dimostra che tra le linee congiungenti due punti P e Q la più corta è quella che si ottiene intersecando la superficie della sfera con il cerchio di centro C e passante per P e Q . Naturalmente dei due archi per P e Q si deve scegliere il più corto è ciò crea un piccolo problema quando i due punti P e Q sono diametralmente opposti. Allora è opportuno considerare solo una parte della superficie sferica in modo che non vi siano punti diametralmente opposti.

Trascurando i particolari tecnici, è immediato rendersi conto che il quinto postulato non vale per tale geometria. Si consideri ad esempio la figura dove i tre punti P , Q ed R sono tali che i piani PCQ e PCR si intersecano in una retta PC ortogonale al piano CQR . Allora nel triangolo PQR gli angoli in Q ed R sono ret-

ti e ciò è in contrasto con il quinto postulato di Euclide. E' anche interessante osservare che il Teorema di Pitagora non vale in tale modello. Basta sempre riferirsi al triangolo con tre angoli retti. In tale triangolo i lati sono uguali, supposto che misurino $a \neq 0$, se valesse il Teorema di Pitagora avremmo che $a^2 = 2a^2$ da cui, dividendo per a^2 avremmo che $1 = 2$.

E facile vedere che in tutti i triangoli la somma degli angoli di tale triangolo è comunque maggiore di un angolo piatto.

Come abbiamo già detto, se si considerano superfici differenti dalla superficie sferica e si definiscono i segmenti come le linee di minima distanza (dette geodetiche), allora si ottengono altri tipi di geometrie. Pertanto esistono tanti tipi di geometrie a due dimensioni quanti sono i tipi di superficie dello spazio. Inoltre, da un punto di vista matematico, non è difficile slittare di una dimensione e considerare una "superficie tridimensionale" immersa in uno spazio a quattro dimensioni. Basta considerare una equazione a quattro incognite, chiamare "spazio quadrimensionale" l'insieme delle sue radici (cioè dei punti di una superficie nello spazio quadrimensionale) e definire al solito i segmenti come le curve continue di tale superficie che siano geodetiche. A tale superficie corrisponderà una geometria a tre dimensioni che in generale risulta essere diversa da quella euclidea.

4. Crisi dell'approccio sintetico: Cartesio

Abbiamo visto come la scoperta delle geometrie non euclidee abbia messo in crisi il convincimento del carattere assoluto della geometria Euclidea. In realtà, prima ancora che con tali scoperte, un primo fondamentale elemento di crisi del metodo di Euclide (se non proprio della sistema geometrico di Euclide) si manifestò con il sistematico processo di algebrizzazione della geometria. Tale processo, iniziato nella prima metà del 1300 con Nicola d'Oresme, trasformerà la geometria "sintetica" di Euclide, in cui si dimostrano teoremi e si tracciano figure, in quella che attualmente si chiama geometria "analitica" in cui tutti i problemi si riducono alla ricerca di radici di sistemi di equazioni algebriche. In un certo senso si passa dalle "dimostrazioni con figure" tipiche della geometria euclidea ai "calcoli" tipici della geometria analitica. Infatti, come è noto, la geometria analitica si ottiene quando, fissati due assi e su di essi due unità di misura, si siano identifica-

- i punti del piano con le relative coordinate,

- le rette con le equazioni di primo grado,
- le coniche con le equazioni di secondo grado e, più in generale, le curve con opportune equazioni implicite o esplicite.

Allora ad ogni operazione geometrica corrisponde una operazione di carattere analitico (cioè relativa ai numeri reali). Ad esempio l'intersezione di due curve si traduce nella risoluzione di un sistema di due equazioni.

Concorsero a tale processo di algebrizzazione scienziati e filosofi come Fermat e Cartesio. In particolare è interessante esaminare il libro di Cartesio *La Geometria* che è una delle appendici del famoso *Discorso sul Metodo* del 1637. La Geometria è costituita da tre parti, di cui la prima porta il titolo "*Dei problemi che si possono costruire col solo uso di cerchi e di linee rette*". Si deve tenere conto che il termine "costruzione" di un problema si deve intendere come "costruzione geometrica di un segmento che sia soluzione del problema" e quindi corrisponde a "risoluzione" di un problema. In questa prima parte si illustra come sia possibile elaborare un "calcolo geometrico" dei segmenti che è l'analogo geometrico della moderna teoria dei numeri reali.

Come l'aritmetica è composta solo di quattro, cinque operazioni, la Somma, la Sottrazione, la Moltiplicazione, La Divisione e la Estrazione delle radici, che si può considerare una specie di Divisione, così anche in Geometria, per quanto riguarda le linee che si cercano . . .

Il brano prosegue spiegando come si possano fare le corrispondenti operazioni con i segmenti. Per la somma e la sottrazione la cosa è evidente. Per quanto riguarda il prodotto e la divisione si utilizza la nozione di proporzione. Infatti supponiamo di volere moltiplicare i segmenti d e c . Allora basta trovare una costruzione geometrica per cui valga una proporzione del tipo $1 : d = c : x$ in quanto, essendo il prodotto dei medi uguale al prodotto degli estremi, in tale caso il segmento x rappresenterà il prodotto di d per c . D'altra parte è ben noto come ottenere grandezze proporzionali in geometria: basta considerare triangoli simili.

Definizione 4.1. Il triangolo ABC si dice *simile* al triangolo $A'B'C'$ se l'angolo in A è uguale all'angolo in A' , l'angolo in B è uguale all'angolo in B' e l'angolo in C è uguale all'angolo in C' .

Ovviamente la relazione di similitudine è una relazione di equivalenza, cioè è riflessiva, simmetrica e transitiva. Ricordando che nella geometria euclidea la somma degli angoli interni di un triangolo è un angolo piatto, è possibile dimostrare la seguente proposizione.

Proposizione 4.2. Dati due triangoli è sufficiente che due degli angoli siano uguali perché si possa asserire che sono simili. Dati due triangoli rettangoli, è sufficiente che uno degli angoli sia uguale per asserire che sono simili.

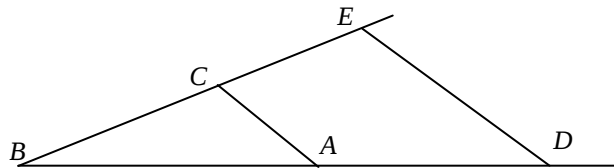
Il seguente teorema, di cui omettiamo la dimostrazione, gioca un ruolo fondamentale nella matematica.

Teorema 4.3. Due triangoli simili hanno i lati proporzionali. Più precisamente, supponiamo che A, B, C siano i vertici di un triangolo e A', B', C' i vertici di un altro triangolo. In tale caso se l'angolo in A è uguale all'angolo in A' , l'angolo in B è uguale all'angolo in B' , e l'angolo in C è uguale all'angolo in C' , allora

$$AB : A'B' = AC : A'C' \quad \text{e} \quad AC : A'C' = BC : B'C'.$$

5. Calcolo dei segmenti

Utilizzando il teorema degli angoli simili Cartesio dice, con riferimento alla seguente figura,



. . . sia per esempio BA l'unità: se bisogna moltiplicare BD per BC devo soltanto aggiungere i punti A e C , poi tracciare DE parallela a CA , e BE è il risultato di questa moltiplicazione.

In altre parole si considerino due rette distinte per il punto B e due punti D e C su tali rette in modo che BD sia uguale a d e BC sia uguale a c . Sia inoltre A un punto della retta per B e D tale

che BA sia unitario. Si tracci infine la parallela a AC per D e si denoti con E il punto di intersezione con la retta per B e C . Allora per la similitudine dei triangoli CBA e EBD risulterà che $1 : BD = BC : BE$. In conclusione BE è il prodotto cercato.

Esercizio. Calcolare il prodotto di 3 per 5 in modo grafico (cioè utilizzando riga e compasso).

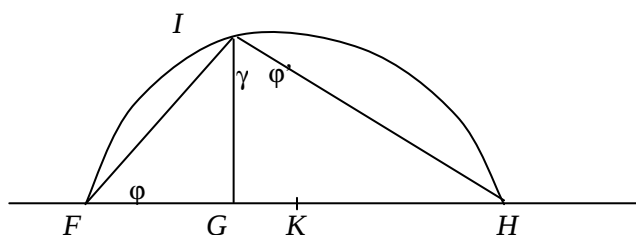
La stessa costruzione, quando si siano dati BE ed BD , vale per la divisione.

. . . Se invece bisogna dividere BE per BD , avendo unito i punti E e D , traccio AC parallela a DE e BC è il prodotto di questa divisione.

Cioè, se si deve dividere il segmento d per il segmento c , allora basta fissare un punto E in modo che BE sia uguale a d ed un punto D in modo che BD sia uguale a c . Tracciata allora la retta per A parallela ad ED , chiamo con C il punto di intersezione con la retta per B ed E . Il segmento BC è il risultato della divisione come si ricava dalla proporzione $1 : BD = BC : BE$.

Esercizio. Calcolare $6/3$ in modo grafico.

Anche nel caso di estrazione di radice quadrata abbiamo che il problema si traduce nello stabilire una opportuna proporzione.



. . . Se bisogna estrarre la radice quadrata di GH , aggiungo in linea retta FG , che è l'unità, e dividendo FH in due parti uguali nel punto K , dal centro K traccio (la semicirconferenza) FIH , poi innalzando dal punto G una linea retta fino ad I ad angoli retti su FH , GI è la radice cercata.

Il fatto che IG sia il segmento cercato è giustificato dal seguente teorema.

Teorema 5.1. La misura del segmento GI è la radice quadrata della misura del segmento GH .

Dim. Da notare che per un teorema della geometria euclidea l'angolo in I è retto.² Per il secondo teorema di Euclide l'altezza di un triangolo rettangolo è media proporzionale delle proiezioni dei cateti. Pertanto vale la seguente proporzione

$$FG : GI = GI : GH.$$

Poiché il prodotto dei medi è uguale al prodotto degli estremi,

$$FG \cdot GH = GI^2$$

e quindi, essendo $FG = 1$, $HG = GI^2$. Se si vuole poi dimostrare tale teorema, è sufficiente mostrare che sono simili i due triangoli FGI e IGH . Ora per mostrare che due triangoli rettangoli sono simili è sufficiente provare che hanno un angolo uguale. Ma ciò non è difficile perché essendo la somma degli angoli di un triangolo un angolo piatto, $\varphi = 180 - 90 - \gamma = 90 - \gamma$ ed essendo il triangolo retto in I , risulta che $\varphi' = 90 - \gamma$. Pertanto $\varphi = \varphi'$. $\square$

Ad esempio se voglio trovare graficamente la radice 9 allora applico la seguente procedura:

1. Traccio un segmento GH di lunghezza 9
 2. Prolungo a sinistra tale segmento con un segmento FG di lunghezza 1
 3. Trovo il punto medio K del segmento FH
 4. Traccio la circonferenza di centro K e diametro $FH = 10$
 5. Alzo la perpendicolare dal punto G
- Il segmento GI misurerà esattamente 3, cioè la radice di 9.

Esercizio. Trovare la radice di 7 in maniera grafica utilizzando cioè un righello ed un compasso.

6. Il “Discorso sul Metodo”

Il calcolo dei segmenti era alla base del metodo proposto da Cartesio, tuttavia si deve sottolineare che quello che Cartesio proponeva era una riduzione della geometria ai suoi elementi più semplici, i segmenti, e non una riduzione della geometria a manipo-

² Il teorema dice che l'angolo al vertice è la metà dell'angolo al centro. Ne segue che il nostro angolo è la metà di un angolo piatto.

lazione algebrica di numeri come avviene attualmente nella geometria analitica.³

Tutti i problemi della geometria si possono facilmente ridurre a tali termini che in seguito per costruirli basta conoscere la lunghezza di alcune rette.

Tali elementi semplici si possono manipolare con operazioni simili a quelle dell'aritmetica e pertanto è più corretto dire che con Cartesio si ha una algebrizzazione della geometria che pone la nozione di operazione alla base di tutto.

D'altra parte in Cartesio non vi era solo l'esigenza di ridurre la geometria a calcolo (di segmenti). Altrettanto importante era il processo inverso che consiste nella possibilità di interpretare ogni discorso algebrico in termini geometrici. In altre parole egli pensava si dovesse

- da un lato liberare la geometria dal ricorso obbligato alle figure che affaticavano inutilmente l'immaginazione
- da un altro lato dare significato alle operazioni dell'algebra per mezzo di una interpretazione geometrica.

Quanto poi all'analisi degli antichi e all'algebra dei moderni . . . , la prima è sempre siffattamente legata alla considerazione delle figure, che essa non può esercitare l'intelligenza senza affaticare di molto l'immaginazione; e, nella seconda, ci si è talmente assoggettati a certe regole e a certe cifre, che se ne è fatta un'arte confusa ed oscura, la quale tiene imbarazzato lo spirito, invece di (essere) una scienza che lo coltivi.

Scopo dichiarato di Cartesio è la ricerca di un metodo generale in contrasto con il modo frammentario con cui procedevano i greci antichi quando si trattava di trovare una dimostrazione o di risolvere un problema. Se infatti è certamente un merito dei greci il fatto che ogni dimostrazione sia rigorosamente controllabile nei suoi singoli passaggi, niente viene detto da essi circa il metodo che si dovrebbe seguire per poter trovare nuovi teoremi e dimostrazioni. Pertanto restiamo disarmati di fronte ad ogni problema

³La geometria analitica intesa come completa riduzione al calcolo numerico non era possibile in quanto nel 1600 non era stata ancora data una definizione di numero reale svincolata dall'intuizione del continuo (bisognerà aspettare per questo la fine del 1800).

nuovo che si presenta e dobbiamo ogni volta procedere per tentativi.

Il metodo proposto da Cartesio per la geometria consisteva

- nell'indicare con lettere i dati e le incognite di un problema geometrico
- di tradurre le informazioni disponibili in equazioni
- nel semplificare, tramite calcoli algebrici, le equazioni quanto più possibile
- nel risolvere le equazioni risultanti da tale semplificazione in termini geometrici.

Pertanto abbiamo un passaggio del tipo

Geometria $\rightarrow$ Algebra $\rightarrow$ Geometria

piuttosto che un annullamento della geometria. Ad esempio dopo aver tradotto un problema geometrico in una equazione di secondo grado era opportuno semplificare al massimo tale equazione. Giunti però alla forma più semplice possibile la risoluzione della equazione finale doveva essere di tipo grafico. Pertanto la risoluzione grafica (detta costruzione) di semplici equazioni di secondo grado, in particolare il calcolo grafico di una radice quadrata, era un elemento essenziale della teoria di Cartesio.

Ecco che cosa dice Cartesio nella sua *Geometria*

Così, volendo risolvere qualsiasi problema, si deve innanzi tutto considerarlo come risolto, e si devono dare dei nomi a tutte le linee che sembrano necessarie per la sua costruzione, sia quelle ignote, sia alle altre. Poi, senza fare alcuna differenza tra queste linee, note ed ignote, bisogna affrontare le difficoltà secondo l'ordine che mostra nella maniera più naturale in che modo tali linee siano in rapporto tra loro, fino a che non si sia trovato modo di esprimere una medesima quantità in due maniere diverse: ciò si chiama un'equazione (in una sola incognita) poiché i termini di una di queste due espressioni sono uguali a quelli dell'altra.

Si noti che Cartesio tratta prevalentemente problemi che si traducono in una equazione ad una sola incognita e che l'idea di luogo geometrico, insieme dei punti le cui coordinate verificano una equazione a due variabili, non è presente nella sua opera se non in modo saltuario.

Concludiamo questo paragrafo sottolineando che le teorie matematiche di Cartesio erano strettamente legate al suo sistema filosofico più generale. Basti pensare che il suo libro *La Geometria* non venne pubblicato come un trattato a sé stante ma come una delle tre appendici del "*Discorso sul metodo*" il cui titolo completo è "*Discorso sul metodo per ben condurre la propria ragione e cercare la verità nelle scienze*" e che tali appendici avevano appunto il ruolo di illustrare il suo metodo filosofico generale. I precetti fondamentali di tale metodo erano:

- Il precetto dell'evidenza;
- Il precetto dell'analisi;
- Il precetto della sintesi;
- Il precetto del computo completo.

Ed il primo era, di non accettare cosa alcuna per vera quando non la riconoscessi evidentemente per tale: cioè, di evitare studiamente la precipitazione e la prevenzione; e di non accogliere nei miei giudizi nulla di più di ciò che si presentasse sì chiaramente e sì distintamente al mio spirito da non poter aver motivo alcuno di metterlo in dubbio.

Il secondo, di dividere ogni difficoltà, ch'io esaminassi, in parti elementari fino al limite del possibile e quanto sarebbe richiesto per trovarne la migliore soluzione.

Il terzo, di condurre per ordine i miei pensieri, cominciando dagli oggetti più semplici e più facili da conoscere, per salire a poco a poco e come per gradi alla conoscenza dei più complessi . . .

E l'ultimo, di fare, in ogni argomento, enumerazioni così complete e verifiche così generali da essere sicuro di nulla omettere.

7. La "costruzione" delle radici di una equazione

Per illustrare il fatto che l'approccio cartesiano non è volto alla sola "traduzione" di ogni problema geometrico in problema algebrico, esaminiamo come ai tempi di Cartesio si usava trasformare un problema algebrico in uno geometrico. Consideriamo ad esempio l'equazione di quarto grado

$$x^4 - x^3 - 3x^2 - 4 = 0. \quad (7.1)$$

Possiamo tentare di abbassare il grado di questa equazione ponendo $y = x^2$ e poi sostituendo al posto di x^2 la variabile y . In tale modo si ottiene il sistema di due equazioni di secondo grado

$$y^2 - xy - 3y - 4 = 0 \quad ; \quad y = x^2.$$

di cui (7.1) è la risultante. Poiché ciascuna equazione può essere vista come l'equazione di una conica, ciò significa che:

è possibile "costruire" le soluzioni di una equazione di quarto grado intersecando due opportune coniche.

Più precisamente dobbiamo disegnare le due coniche, intersecarle e poi andare a vedere le ascisse dei punti di intersezione. È il processo inverso a quello a cui siamo abituati: quello per cui, dovendo trovare i punti di intersezione di due coniche, scriviamo il sistema delle relative equazioni, troviamo l'equazione risultante e poi risolviamo tale equazione con qualche formula. Poiché la risoluzione di una equazione di quarto grado è un problema complicato, lo si traduce nel più semplice problema (grafico) di intersecare due coniche. In questo modo si capisce anche perché, in generale, una equazione di quarto grado ha quattro soluzioni.

Naturalmente se invece della semplice sostituzione $y = x^2$ si utilizza qualche parametro, allora è possibile ottenere altri tipi di coniche. Ad esempio se si pone

$$y = x^2 + \lambda x + \mu$$

allora elevando al quadrato entrambi i membri dell'equazione si ottiene

$$y^2 = x^4 + \lambda^2 x^2 + \mu^2 + 2\lambda x^3 + 2\mu x^2 + 2\lambda\mu x.$$

Da tale equazione si ricava che $x^4 = y^2 - \lambda^2 x^2 - \mu^2 - 2\lambda x^3 - 2\mu x^2 - 2\lambda\mu x$ e sostituendo in (7.1) si ottiene $x^4 - x^3 - 3x^2 - 4 = 0$.

$$y^2 - (2\lambda + 1)x^3 - (2\mu + \lambda^2 + 3)x^2 - 2\lambda\mu x - \mu^2 - 4 = 0.$$

Ponendo $\lambda = -1/2$ eliminiamo x^3 , ed otteniamo

$$y^2 - (2\mu + 13/4)x^2 + \mu x - \mu^2 - 4 = 0.$$

Possiamo ora scegliere il parametro μ in modo opportuno. Ad esempio se scegliamo μ in modo che $2\mu + 13/4 = 0$, cioè $\mu = -13/8$ tale equazione diventa l'equazione di una parabola. Se scegliamo μ in modo che $2\mu + 13/4 = -1$, tale equazione diventa quella di un cerchio. In definitiva possiamo ottenere le radici dell'equazione (7.1) sia intersecando due parabole, sia intersecando una parabola con un cerchio.

Quanto fatto per l'equazione (7.1) può essere fatto per ogni equazione di quarto grado.

Proposizione 7.1. Le radici di una equazione di quarto grado possono essere trovate considerando le ascisse dei punti di intersezione di una parabola fissata e di una iperbole che dipende dalla equazione data.

Dim. Data l'equazione $x^4+ax^3+bx^2+cx+d=0$, poniamo $y=x^2$. Si ottiene $y^2+axy+by+cx+d=0$ e quindi l'equazione di quarto grado è la risultante dell'equazione $y=x^2$ di una parabola e dell'equazione $y(y+ax)+by+cx+d=0$ di una iperbole.

Nel caso di equazione di terzo grado le cose sono ancora più semplici. Ad esempio consideriamo l'equazione

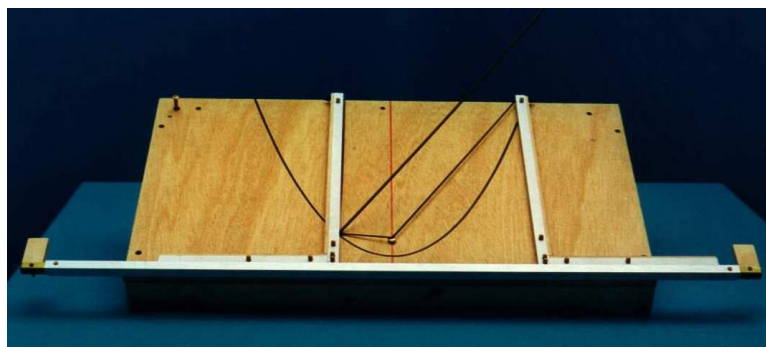
$$x^3-2x^2+x-1=0.$$

Posto $y=x^2$ si ottiene che tale equazione è la risultante delle equazioni

$$xy-2y+x-1=0 \quad ; \quad y=x^2,$$

che rappresentano una iperbole ed una parabola.

Da notare che il disegno delle coniche veniva fatto con “macchine” che tracciavano le curve in modo meccanico allo stesso modo come un compasso traccia un cerchio. Ad esempio in un bel sito sulle “macchine matematiche”⁴ viene riportata la seguente macchina per tracciare parabole:



⁴ Si veda http://www.museo.unimo.it/theatrum/macchine/_00lab.htm.

8. Aritmetizzazione della geometria: la sparizione delle figure

*“Il lettore non troverà figure in questo lavoro. I metodi che esporrò non richiedono costruzioni né geometriche né meccaniche, ma solamente operazioni algebriche, soggette a una procedura regolare e uniforme.” Jean Louis Lagrange, **Mécanique Analytique***

Se si considera l'importanza che la geometria di Euclide aveva avuto nella cultura dei greci ed in quella successiva, ci si rende conto di quanto fosse un evento rivoluzionario e sorprendente la scoperta delle geometrie non euclidee. I pensatori precedenti avevano costantemente ritenuto che vi fosse una sola geometria vera e che le sue leggi fossero necessariamente quelle di Euclide. Inoltre il modo di procedere geometrico era sempre stato visto come un modello a cui ispirarsi in tutti gli altri campi del sapere. L'apparire di tali nuove geometrie confutava queste convinzioni perché se più teorie dello spazio contrastanti tra loro sono logicamente possibili e se solo una di queste poteva essere vera, allora la geometria, e più in generale la matematica non poteva più essere considerata lo strumento per giungere alla verità.

D'altra parte il mettere in discussione il ruolo centrale ed assoluto della geometria faceva nascere in modo sempre più pressante l'esigenza di trovare una nuova base alla matematica. Si deve anche tenere conto che era necessario inquadrare anche la nuova matematica nata dai metodi infinitari del calcolo differenziale.

I passi di una tale nuova fondazione della matematica consistono essenzialmente:

- a) nell'aritmetizzazione della geometria e dell'analisi,
- b) nella teoria degli insiemi di G. Cantor.

Abbiamo già visto che con Cartesio la geometria era stata ridotta ad un calcolo dei segmenti e quindi, in un certo senso, ridotta all'algebra. Per poter fare completamente a meno della geometria era allora necessario un ulteriore passo: sostituire al calcolo dei segmenti un calcolo numerico che si fondasse su di una definizione di numero reale completamente indipendente dall'intuizione geometrica.

Nel prossimo capitolo mostreremo come ciò sia possibile definendo prima i numeri naturali, poi gli interi relativi, poi i razionali ed infine i reali. In questo paragrafo daremo per scontata la conoscenza dei numeri reali e mostreremo come tali numeri siano sufficienti a definire le nozioni geometriche. Infatti, come è noto, con lo sviluppo della geometria analitica, la geometria diverrà un capitolo dell'algebra lineare sul campo dei numeri reali.

Teorema 8.1. Sia R l'insieme dei numeri reali, chiamiamo *piano euclideo* il prodotto cartesiano $R \times R$, e *punti* i suoi elementi. Inoltre chiamiamo *retta* un insieme di punti che verifichi una equazione lineare del tipo $ax+by+c=0$. Allora la struttura ottenuta in tale modo verifica tutti gli assiomi della geometria euclidea.

Dim. Proviamo ad esempio che per due punti distinti (x_0, y_0) e (x_1, y_1) passa una ed una sola retta. Tale problema si traduce in quello di trovare a, b, c (non tutti nulli) tali che

$$ax_0+by_0+c=0$$

$$ax_1+by_1+c=0$$

Si tratta di un sistema omogeneo di due equazioni nelle incognite a, b e c e dalla teoria dei sistemi lineari si sa che se i due punti sono diversi tra loro tale sistema ammette infinite soluzioni e che due diverse soluzioni sono proporzionali tra loro. Ne segue che tutte queste soluzioni rappresentano una stessa retta e ciò prova l'esistenza e l'unicità della retta per (x_0, y_0) e (x_1, y_1) . Più in generale, possiamo dire che il passaggio per un punto è una condizione lineare omogenea e che quindi la retta che passa per due punti prefissati distinti si trova imponendo due condizioni lineari omogenee. Da ciò segue l'esistenza e l'unicità di tale retta.

In modo analogo possiamo provare poi il quinto postulato. Si tratta di verificare che dato un punto $P = (x, y)$ ed una retta r , che supponiamo di equazione $ax+by+c=0$, allora esiste una ed una sola retta passante per (x, y) e parallela ad r . Ma anche il parallelismo è una condizione lineare omogenea poiché una retta r' con coefficienti a' e b' è parallela ad r se e solo se $ab'-a'b=0$. Pertanto esiste ed è unica la retta per P parallela ad r . In modo altrettanto semplice si dimostrano i rimanenti assiomi. $\square$

Un processo di riduzione al calcolo dei numeri reali è stato poi fatto anche per quanto riguarda l'analisi matematica. Infatti nascono le attuali definizioni di limite, derivata, funzione continua

che fino alla prima metà dell'ottocento si basavano sulla intuizione del continuo geometrico. In definitiva si attua quello che viene a volte chiamato “*processo di aritmetizzazione della matematica*” in cui tutto viene ridotto ai numeri reali. Poiché i numeri reali si possono definire a partire dagli interi, riappare l'antica idea della scuola Pitagorica per cui tutto è riconducibile ai numeri interi.

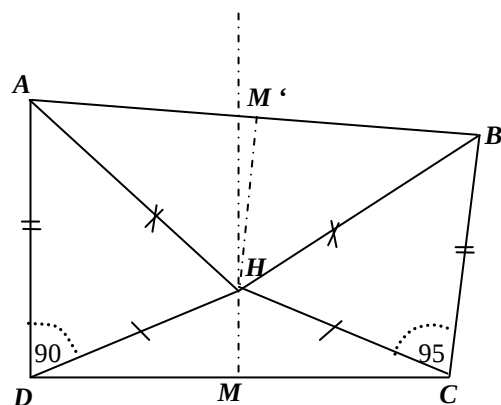
9. Intuizione geometrica e falsi teoremi euclidei

Le dimostrazioni presenti nei libri di Euclide sono sempre molto belle ed intuitive. Infatti in esse sono sempre presenti sia il rigore logico della deduzione sia l'interpretazione intuitiva dei singoli passi di tale deduzione. Tuttavia a volte l'intuizione geometrica trae in inganno (come abbiamo già visto nel caso della equiscomponibilità) ed in questo paragrafo mostro alcuni esempi. Nel seguito riporto la dimostrazione, di tipo euclideo, del fatto che $5 = 0$. Non so chi l'abbia inventata in quanto mi è stata raccontata da un collega a cui è stata raccontata da un altro collega . . . Lascio a chi legge il compito non facile di trovare dove è l'errore.

Teorema 9.1. Il numero 5 è uguale al numero 0.

Dim. Tracciamo un segmento DC ed alziamo da D un segmento AD perpendicolare a DC . Alziamo da C un segmento di uguale lunghezza che faccia un angolo di 95 gradi col segmento CD . Otteniamo un quadrilatero di vertici A, B, C, D . Tracciamo ora l'asse di DC (dal punto medio M di DC) e l'asse di AB (dal punto medio M' di AB). Poiché DC non è parallelo ad AB i due assi non sono paralleli tra di loro e pertanto si incontrano in un punto H . Si vengono pertanto a formare due triangoli AHD e BHC che risultano uguali. Infatti AD è uguale a BC per costruzione, $AH = HB$ perché il triangolo AHB è isoscele, $DH = CH$ perché il triangolo DHC è isoscele. Dal fatto che AHD sia uguale a BHC comporta che l'angolo $\hat{ADH}$ sia uguale all'angolo $\hat{HCB}$. Essendo il triangolo DHC isoscele, risulta anche che $\hat{HDC} = \hat{DCH}$. In definitiva possiamo concludere che

$$\hat{ADC} = \hat{ADH} + \hat{HDC} = \hat{HCB} + \hat{HDC} = \hat{DCB}$$



e quindi che $90 = 95$. Sottraendo 90 si ottiene che $0 = 5$. $\square$

Corollario 9.2. Io sono l'uomo più bello, più intelligente, più simpatico del mondo, inoltre sono anche il più grande matematico che sia mai esistito.

Dim. Se consideriamo l'insieme costituito da me, dall'uomo più bello, dall'uomo più intelligente, dall'uomo più simpatico e dal più grande matematico che sia mai esistito arriviamo ad un insieme X con 5 elementi. Poiché abbiamo dimostrato che $5 = 1$, X ha un solo elemento e quindi tutte le persone che ho elencato prima sono in realtà una unica persona. Ciò prova il corollario⁵.

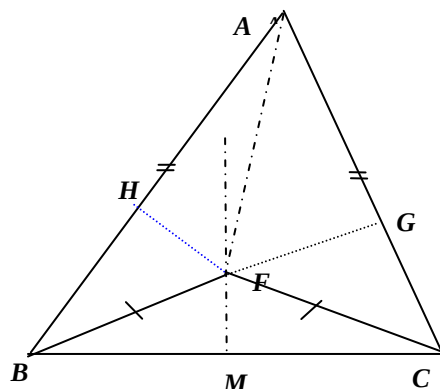
Un altro falso teorema (in cui viene fatto un errore simile a quello del teorema ora esposto) è il seguente.

Teorema 9.3. Tutti i triangoli sono isosceli.

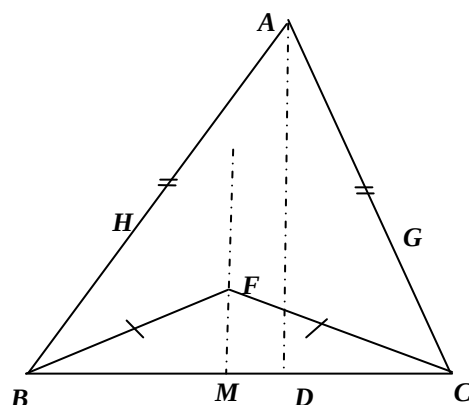
Dim. Consideriamo un triangolo di vertici A, B, C i suoi vertici e con M il punto medio del segmento BC . Alziamo da M la perpen-

⁵ Questa dimostrazione è un ovvio adattamento di una dimostrazione che si racconta abbia fatto Bertrand Russell. Pare a Russell sia stato chiesto come possa essere accettata una cosa tanto strana per cui a partire da una asserzione falsa possa essere dimostrata qualsiasi altra asserzione e che per sfida gli sia stato chiesto come da $2=1$ si possa dimostrare, ad esempio, che Russell è Dio. La risposta di Russell fu appunto che essendo Russell e Dio due cose distinte ed essendo $2=1$, allora Russell non poteva che essere uguale a Dio.

dicolare a BC e da A la bisettrice dell'angolo in A . Due sono i casi: che le due rette si incontrino in un punto o che siano parallele (si vedano le figure affianco). Nel primo caso indichiamo con F il punto di incontro e, dopo avere tracciato i segmenti FB ed FC , a partire da F tracciamo le perpendicolari FG e FH ad AC e AB . I due triangoli AHF e AFG sono uguali in quanto sono rettangoli, hanno un lato in comune e, essendo la retta AF la bisettrice, l'angolo HAF è uguale all'angolo FAG . D'altra parte i due triangoli BFM e MFC sono uguali in quanto rettangoli con $BM = MC$ ed il lato MF in comune. Ne segue che essendo $BF = FC$ e $HF = FG$ i due triangoli rettangoli BHF e FCG sono uguali. Concludendo $AB = BH + HA = CG + AG = AC$ e quindi ABC è isoscele.



Consideriamo ora il caso in cui la bisettrice e l'asse BC non si incontrano.



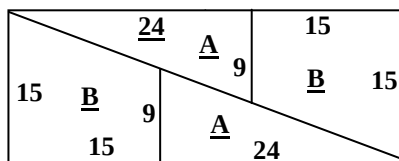
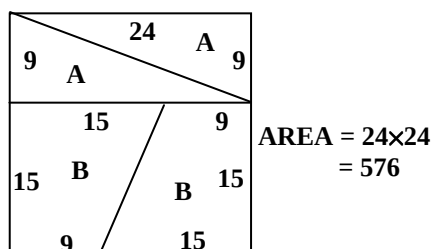
Allora le due rette sono parallele e quindi la bisettrice risulta perpendicolare a BC . Ne segue che i due triangoli rettangoli BDA e

ADC avendo un lato in comune ed i due angoli in A uguali sono uguali. Pertanto $BA = AC$. $\square$

Ora mostro un esempio di coppia di figure che sembrerebbero equiscomponibili ma non lo sono.

Proposizione 9.4. (Paradosso della scomposizione del quadrato). E' possibile scomporre un quadrato di area 576 ed ottenere un rettangolo di area 585.

Dim. Costruiamo un quadrato di lato uguale a 24 e tagliamolo in due rettangoli di lati 24 e 9 e 24 e 15. Tagliamo il rettangolo piccolo in due triangoli rettangoli di cateti 9 e 24 ed il rettangolo grande in due trapezi rettangoli di cateti 15 e 9 e di ipotenusa uguale a 15. Ricomponiamo poi tali pezzi in modo da ottenere il rettangolo dato in figura di lati 15 e 39 (si vedano le figure). Per costruzione il quadrato ed il rettangolo sono equiscomponibili.



$$AREA = (15+24) \times 15 = 39 \times 15 = 585$$

Inoltre l'area del quadrato è uguale a $24 \times 24 = 576$, l'area del rettangolo risulta uguale a $15 \times 39 = 585$.⁶ $\square$

⁶ Chi non fosse convinto dei disegni può provare a fare esperimenti con un foglio di carta quadrettato ed un paio di forbici.

CAPITOLO 3

DEFINIRE I NUMERI

*Dio creò i numeri naturali, tutto
il resto è opera dell'uomo.*
Leopold Kronecker (1823-1891)

1. Un punto di partenza: terne di Peano

In questo capitolo vogliamo mostrare come sia possibile definire i numeri, ed in particolare i numeri reali i quali, come abbiamo già osservato nel capitolo precedente, possono essere considerati la base su cui fondare una buona parte della matematica.¹ Cominciamo con i numeri naturali, cioè gli interi positivi 0, 1, 2, ... Tali numeri costituiscono una nozione tanto immediata che probabilmente pretendere di definirla, quindi di ridurla a termini più semplici, non ha molto senso. Tuttavia lo sforzo di definirli ha il vantaggio di mettere in rilievo le proprietà essenziali di tali numeri. Il sistema comunemente accettato è quello dovuto a Dedekind ed a Peano che assiomatizzano l'idea intuitiva per cui l'insieme dei numeri naturali è il frutto del “processo di aggiungere un nuovo elemento ad un elemento dato” e quindi di “successivo” (si veda il raccontino di Zavattini alla fine del capitolo).

Definizione 1.1. Consideriamo una struttura algebrica (S, s, z_0) con s operazione 1-aria ed $z_0 \in S$ elemento designato. Diciamo che tale struttura è una *terna di Peano* se sono verificati i seguenti assiomi:

P1 $s : S \rightarrow S$ è una funzione iniettiva

P2 $z_0 \notin s(S)$, cioè z_0 non è il successivo di nessun elemento

P3 per ogni sottoinsieme D di S

$$z_0 \in D \text{ e } s(D) \subseteq D \Rightarrow D = S.^2$$

¹ A tale scopo utilizzeremo alcune nozioni elementari di teoria degli insiemi e relative alle strutture algebriche, argomenti che supporremo già noti al lettore ma che comunque saranno esposti nei prossimi capitoli.

² Da notare che tale teoria è espressa “al secondo ordine”. Ciò significa che si utilizza un linguaggio in cui si applica un quantificatore (“per ogni”) a sottoinsiemi D di S . In logica matematica, come vedremo, normalmente si considerano invece teorie “del primo ordine” in cui è possibile quantificare solo su elementi di S . Esiste pertanto anche una teoria del primo ordine delle terne di Peano che vedremo nel seguito.

La funzione $s : S \rightarrow S$ viene chiamata *funzione-successore*, l'elemento z_0 viene chiamato *elemento nullo*, $s(x)$ il *successivo* di x . L'assioma **P3** è quello più importante e prende il nome di *principio di induzione matematica*. **P3** può anche essere scritto al modo seguente:

$$((z_0 \in X) \text{ e } (x \in X \Rightarrow s(x) \in X)) \Rightarrow X = S.$$

Quando si propone una teoria si deve anche provare che ne esiste almeno un modello. In caso contrario la teoria parlerebbe del nulla. Pertanto dovremmo provare che esiste almeno una terna di Peano, cosa che faremo in seguito. Non è difficile trovare comunque esempi “concreti” di terne di Peano, per meglio dire esempi di esperienze da cui sia possibile far nascere la nozione di terna di Peano tramite un opportuno processo di astrazione. Ne esponiamo due.

Terne di Peano e tacche di legno: Sicuramente uno dei sistemi utilizzati dagli uomini primitivi per contare le pecore di un gregge è quello di mettere delle tacche su di un pezzo di legno o su di un osso. Questo suggerisce che l'insieme delle possibili tacche su un pezzo di legno costituisce un esempio di terna di Peano. In tale caso un pezzo di legno senza tacche rappresenta il primo elemento, l'operazione di aggiungere una tacca è l'operazione successore. Per maggiore precisione dobbiamo identificare due pezzi di legno che abbiano la stessa quantità di tacche come rappresentativi dello stesso numero.

Terne di Peano e sassolini: Un altro mezzo per contare è quello di utilizzare sassolini. Consideriamo ad esempio recipienti contenenti sassolini. Un recipiente vuoto corrisponde allo zero. L'operazione di aggiungere un sassolino ad un recipiente corrisponde all'operazione di successivo.

Naturalmente tali esempi non sono di tipo matematico e se volessimo essere più rigorosi dovremmo procedere a qualche forma di “idealizzazione”. Ad esempio nel caso dei pezzi di legno con tacche si deve immaginare che esistano infiniti possibili pezzi di legno, almeno uno per ogni possibile sequenza di tacche. Inoltre se due pezzi di legno hanno tre tacche, allora devono essere considerati equivalenti, cioè rappresentativi di un solo “oggetto ideale” (il numero 3).

E' anche interessante far vedere che molte strutture che i matematici utilizzano usualmente non sono terne di Peano: ecco alcuni esempi.³

Sia Z l'insieme degli interi relativi, allora $(Z, s, 0)$, dove $s(x) = x+1$ non è una terna di Peano. Infatti l'assioma **P2** non vale in quanto 0 è il successore di -1 . Inoltre non vale nemmeno **P3**. Infatti l'insieme D degli interi maggiori o uguali a zero pur verificando le due condizioni $0 \in D$ e $D \subseteq s(D)$ non coincide con Z .

Sia R^+ l'insieme dei reali maggiori o uguali a zero, allora $(R^+, s, 0)$ non è una terna di Peano. Infatti anche se gli assiomi **P1** e **P2** sono soddisfatti **P3** non è soddisfatto.

Sia Z/m l'insieme degli interi modulo m e consideriamo la struttura $(Z/m, s, [0])$ dove $s([x]) = [x] + [1] = [x+1]$. E' evidente che **P3** è verificata, tuttavia **P2** non vale in quanto $[0]$ è successore di $[m-1]$. Questo mostra che $(Z/m, s, [0])$ non è una terna di Peano.

Problema: Esiste una terna di Peano con 5 elementi ?

Problema. Dimostrare che $(R^+, s, 0)$ non è una terna di Peano fornendo almeno due esempi di insieme per cui non vale il principio di induzione.

Problema. Consideriamo la struttura (N_0, s, z_0) con N_0 insieme dei numeri naturali, $z_0 = 0$ ed s definita dal porre $s(n) = 2n+1$. Dire se tale struttura è una terna di Peano.

Esercizio. Consideriamo la struttura (D, s, z_0) con D insieme dei numeri naturali dispari, $z_0 = 0$ ed s definita dal porre $s(n) = n+2$. Dire se tale struttura è una terna di Peano.

Esercizio. Consideriamo la struttura (S, s, z_0) con S insieme dei numeri naturali maggiori o uguali a 5, $z_0 = 5$ e porre $s(n) = n+1$. Dire se tale struttura è una terna di Peano.

Esercizio. Sappiamo che il prodotto diretto di due gruppi è un gruppo e lo stesso si può dire per gli anelli o per i reticoli. Consideriamo il prodotto diretto di una terna di Peano (P, s, z_0) per se

³ Dare tali esempi sarebbe scorretto da un punto di vista metodologico. Infatti se stiamo "fondando" la matematica non possiamo riferirci ad esempi che attingono da una matematica non ancora fondata. La scorrettezza è solo apparente in quanto il ruolo di tali esempi è didattico (attingere ad una intuizione già esistente) e non matematico. D'altra parte che senso avrebbe l'impresa di "fondare la matematica" se non si avesse in mente già una idea dell'oggetto-matematica ?

stessa, cioè la struttura $(P \times P, s, (z_0, z_0))$ definita ponendo $s((x, y)) = (s(x), s(y))$. Dire se tale struttura è ancora una terna di Peano.

2. Principio di induzione

In una terna di Peano è possibile effettuare due cose di particolare importanza: le dimostrazioni tramite il principio di induzione e le definizioni per ricorsione.

Proposizione 2.1 (Principio di induzione⁴). Supponiamo che una proprietà P sia definita in una terna di Peano (S, s, z_0) e che P verifichi le due seguenti affermazioni:

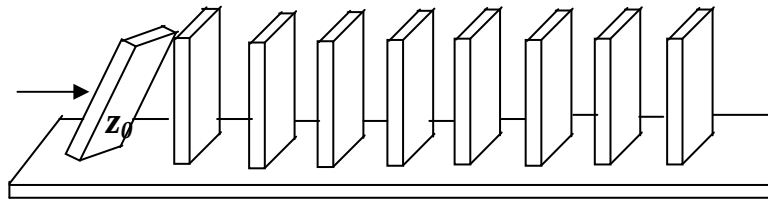
1. P è verificata da z_0
2. se P è verificata da x allora è verificata da $s(x)$

allora è possibile concludere che

P è verificata per ogni $x \in S$.

Dim. Sia D l'insieme degli elementi che verificano P , allora D contiene z_0 ed è tale che $s(D) \subseteq D$. Pertanto tale insieme coincide con S e ciò prova che P è verificata per ogni $x \in S$. $\square$

L'applicazione del principio di induzione può essere visualizzata al modo seguente. Consideriamo la seguente figura in cui i pezzi del gioco domino sono poggiati su di un tavolo (infinito) uno dopo l'altro:



Indichiamo il primo pezzo della fila con z_0 . Vale la regola che se un pezzo cade (a destra) allora il pezzo successivo cade.

⁴ Non bisogna confondere tale principio, che appartiene alla matematica, con il principio di induzione in fisica. In fisica il principio di induzione è quello che permette di passare da una serie di esperimenti alla formulazione di una teoria. Ad esempio, poiché in tutte le nostre esperienze passate un corpo libero cade verso la terra, possiamo formulare la teoria che per ogni corpo x , se x è libero allora x cadrà verso la terra.

$$\forall x \text{Cade}(x) \rightarrow \text{Cade}(s(x))$$

poi supponiamo che valga

$$\text{Cade}(z_0)$$

Allora vale $\forall x \text{Cade}(x)$, cioè tutti i pezzi cadono.

Esempio. Dimostriamo per induzione che per ogni $n \in \mathbb{N}$, $2^n > 0$. Infatti la diseuguaglianza è vera per $n = 0$. Supposta vera per n , supposto cioè che $2^n > 0$, risulta che $2 \cdot 2^n > 2 \cdot 0 = 0$ e quindi che la diseuguaglianza è vera per $n+1$. Pertanto la diseuguaglianza vale per ogni $n \in \mathbb{N}$.

Esempio. Supponiamo di voler dimostrare la seguente asserzione

“la somma dei primi n numeri è uguale a $n \cdot (n+1)/2$ ”

Per $n = 1$ l'asserzione è vera. Supponiamo che l'asserzione sia vera per n . Allora la somma dei primi $n+1$ numeri è $n \cdot (n+1)/2 + (n+1) = (n \cdot (n+1) + 2 \cdot (n+1))/2 = (n+1) \cdot (n+2)/2$. In definitiva se l'asserzione è vera per n è vera anche per $n+1$. Per il principio di induzione l'asserzione è vera per ogni n .

Falsa dimostrazione: Trovare l'errore nella dimostrazione del seguente teorema:

Teorema: Tutte le persone hanno la stessa età.⁵

Dim: Indichiamo con $S(n)$ l'asserzione “in un gruppo di n persone tutte hanno la stessa età”.

Passo 1: $S(1)$ è ovviamente vera

Passo 2: Supponiamo che $S(n)$ sia vera: vogliamo provare che $S(n+1)$ è vera. Sia G un gruppo con $n+1$ persone e siano P_1 e P_2 due persone del gruppo. Allora, detta P una qualunque persona diversa da P_1 e P_2 in $G - \{P\}$ esistono n persone e quindi, per ipotesi di induzione, in $G - \{P\}$ tutte le persone hanno la stessa età. In particolare P_1 ha la stessa età di P_2 .

⁵ Questa falsa dimostrazione ed il successivo paradosso considerano come punto di partenza 1 e non 0 come richiederebbe il principio di induzione. Questo fatto è giustificato da una ovvia estensione di tale principio considerata nella proposizione 6.1.

Paradosso:

Teorema: Se si accetta che non tutti i mucchi di grano sono piccoli allora il principio di induzione è falso.

Dim: E' ovviamente vero che un mucchio di grano costituito da un solo chicco è piccolo. Inoltre se un mucchio di grano è piccolo, allora rimane piccolo anche se ci aggiungo un chicco di grano. Se indichiamo con $P(n)$ l'asserzione "un mucchio con n chicchi è piccolo" queste due verità possono essere rappresentate brevemente al modo seguente

1. $P(1)$ è vera

2. per ogni $n \in \mathbb{N}$, $P(n) \rightarrow P(n+1)$ è vera.

Se il principio di induzione fosse valido, allora $P(n)$ dovrebbe essere vera per ogni n in contrasto con le ipotesi.

3. Definizione per ricorsione

Il principio di induzione permette di definire "per ricorsione" funzioni ed operazioni sui numeri naturali. Ad esempio consideriamo la funzione fattoriale che indichiamo con $fatt$. Di solito tale funzione viene definita dicendo che, per ogni naturale n , $fatt(n)$ è il prodotto dei primi n numeri, oppure si scrive $fatt(n) = 1 \cdot 2 \cdot \dots \cdot n$. Un modo più elegante e preciso di definire il fattoriale è dire che è l'unica funzione che soddisfa le condizioni

$$fatt(0) = 1 \quad ; \quad fatt(n+1) = fatt(n) \cdot (n+1).$$

Per fare un altro esempio consideriamo il seguente problema.

Il problema delle strette di mano. Poniamoci il seguente problema:

Quante strette di mano devono darsi 3 persone che si incontrano in una festa ?

La risposta è semplice e diretta ed è il numero 3 (3 è abbastanza piccolo da permetterci di immaginare direttamente la scena delle strette di mano). Passiamo ora alla domanda:

Quante strette di mano devono darsi 7 persone che si incontrano in una festa ?

Si invita chi legge a dedicare un po' di tempo a risolvere questo problema. Si accorgerà che la risposta richiede un minimo di tempo e pazienza. La risposta diventa poi difficile al posto di 7 si considera un numero più grande, ad esempio il numero 10.

Paradossalmente è più semplice invece affrontare il problema in generale e chiedersi:

Quante strette di mano devono darsi n persone che si incontrano in una festa?

Se indichiamo con $f(n)$ tale numero possiamo tentare di calcolare i primi valori della funzione f . E' evidente che $f(1) = 0$. Infatti in un gruppo con una sola persona non possono esserci strette di mani. E' anche facile vedere che $f(2) = 1$ e che, come abbiamo già visto, $f(3) = 3$ ma già il calcolo di $f(4)$ si presenta un po' noioso. . .

Tuttavia un buon matematico dovrebbe accorgersi che l'avere calcolato $f(3)$ può essere utilizzato nel calcolo di $f(4)$. Infatti se dopo che tre persone si sono strette la mano arriva alla festa una nuova persona a tale nuova persona non resta che fare tre strette di mano. Quindi $f(4) = f(3)+3 = 6$. Se poi alle 4 persone si aggiunge un nuovo venuto, allora si devono aggiungere ancora quattro strette a quelle già fatte. Pertanto $f(5) = f(4)+4 = 10$. Più in generale, per rispondere alla nostra domanda basta calcolare i valori della successione

$$\begin{aligned} f(1) &= 0, f(2) = f(1)+1 = 1, f(3) = f(2)+2 = 3, f(4) = f(3)+3 = 6, \\ f(5) &= f(4)+4 = 10, f(6) = f(5)+5 = 15, f(7) = f(6)+6 = 21, \\ f(8) &= f(7)+7 = 28, f(9) = f(8)+8 = 36, f(10) = f(9)+9 = 45. \end{aligned}$$

Si osservi che la funzione f è completamente definita dalle due equazioni

$$f(1) = 0 ; f(n+1) = f(n)+n$$

che mostrano come, similmente a quanto avviene per il fattoriale, si possa calcolare il valore di f in un numero in funzione del valore di f in un numero precedente.⁶ Questi due esempi suggeriscono la seguente definizione.

Definizione 3.1. Sia (S, s, z_0) una terna di Peano ed $f : S \rightarrow S$ una funzione. Diciamo che f è definita per *ricorsione* tramite l'elemento $c \in S$ e la funzione $h : S^2 \rightarrow S$ se soddisfa le equazioni:

$$f(z_0) = c ; f(s(n)) = h(n, f(n)). \quad (3.1)$$

⁶ E' facile vedere che $f(n) = 1+2+\dots+n-1$, cioè che $f(n)$ è la somma dei primi $n-1$ numeri naturali. Abbiamo già incontrato una funzione simile quando abbiamo parlato dei numeri triangolari. Si prova, per induzione su n , che $f(n) = n \cdot (n-1)/2$.

La prima equazione in (3.1) viene detta "*assegnazione iniziale*" mentre la seconda "*schema di ricorsione*". Lo schema di ricorsione dice che il valore di f in un numero può essere calcolato in funzione del valore di f nel precedente di tale numero. Nello schema di ricorsione compare la stranezza per cui si definisce una cosa utilizzando la cosa stessa, infatti si definisce f "ricorrendo" ad f stesso. Tuttavia si deve osservare che la funzione f a destra dello schema viene applicata ad un numero n che è più semplice del numero $s(n)$ che compare a sinistra. Pertanto applicando più volte lo schema si finisce con il dovere applicare f al valore z_0 , cosa questa che viene consentita dall'assegnazione iniziale.

Fissati c ed h possiamo vedere le due equazioni in (3.1) come una "definizione" della funzione f . Tuttavia è necessario stare attenti all'uso dell'espressione "definizione". Infatti quando definiamo un ente matematico tramite una serie di proprietà allora la definizione è corretta solo se esiste un ed un solo ente verificante tale proprietà. Ad esempio se, nell'ambito della teoria dei numeri reali dico "sia r la radice di -2 " ho una definizione non corretta poiché non esiste nessuna numero reale che soddisfa tale condizione. Se dico "sia r la radice di 2 " ho una definizione non corretta poiché esistono due numeri reali che soddisfano tale condizioni. Una definizione corretta è invece, ad esempio, "sia r la radice positiva di 2 ". Ovviamente anche per la definizione 3.1 si pone la stessa questione. Da notare che in ambiti che non sono le terne di Peano tale teorema di esistenza e di unicità non è detto che valga. Ad esempio vale la seguente proposizione.

Proposizione 3.2. Nell'anello degli interi modulo m non esiste nessuna funzione $fatt$ che verifica le equazioni

$$fatt(0) = 1; fatt(x+1) = fatt(x) \cdot (x+1) \quad (3.2).$$

Nell'insieme R^+ dei numeri reali maggiori o uguali a 0 esistono infinite funzioni che verificano tali equazioni.

Dim. Se negli interi modulo m esistesse una tale funzione allora dovrebbe essere

$$fatt(0) = fatt(m) = fatt(m-1) \cdot m = fatt(m-1) \cdot 0 = 0$$

in contrasto con la condizione $fatt(0) = 1$.

Consideriamo le stesse equazioni in R^+ , allora è subito visto che le due equazioni impongono condizioni solo sui numeri naturali. Questo significa che tutte le funzioni il cui grafico passa per

i punti di coordinate $(0,1), (1,1), (2,2), (3,6), \dots, (n, fatt(n)), \dots$ verificano le equazioni (3.2). $\square$

Nelle terne di Peano invece le cose funzionano bene.

Teorema 3.3. Sia (S, s, z_0) una terna di Peano, allora per ogni $c \in S$ ed $h : S^2 \rightarrow S$ esiste ed è unica una funzione f definita in tutto S che soddisfa le equazioni in (3.1). Tale funzione è totale.

Dim. Per quando riguarda il problema dell'esistenza nella maggior parte dei testi tale esistenza viene data come fatto ovvio (in generale nei testi di informatica). Infatti (3.1) rappresenta un algoritmo per effettuare un calcolo e tale algoritmo fornisce un output per ogni possibile input. D'altra parte sembra naturale accettare che se esiste un algoritmo esiste anche la funzione corrispondente. Abbiamo visto tuttavia che tale esistenza non è verificata se ci si riferisce agli interi modulo m e che quindi le cose sono meno semplici di come appare. Pertanto in molti testi si sente il bisogno di effettuare una dimostrazione di tale esistenza (in generale nei testi di algebra). Accenno solo a come si dovrebbe procedere. Si considera la classe C di tutte le relazioni binarie $\mathcal{R}$ in S che soddisfano la seconda equazione in (3.1), cioè tali che

$$(n, m) \in \mathcal{R} \Rightarrow (s(n), h(n, m)) \in \mathcal{R}. \quad (3.3)$$

Tale classe è non vuota in quanto la relazione totale $\mathcal{R} = S \times S$ verifica (3.3). Inoltre C costituisce un sistema di chiusura e quindi possiamo considerare la relazione $\mathcal{R}$ generata dalla coppia (z_0, c) .⁷ Si prova (cosa alquanto noiosa) che tale relazione è una funzione ovunque definita che è, appunto, la funzione di cui si vuole provare l'esistenza.

Per provare esplicitamente che tale funzione è ovunque definita basta osservare che per la prima equazione f è definita in z_0 e che, per la seconda equazione, se f è definita in n allora è definita anche in $s(n)$. Per il principio di induzione ciò significa che f è definita per ogni n . Per provare l'unicità supponiamo che f ed f' siano due funzioni soddisfacenti le equazioni (3.1) e sia $X = \{x \in S : f(x) = f'(x)\}$. Allora è evidente che $z_0 \in X$ e che se $n \in X$ allora, poichè

$$f(s(n)) = h(n, f(n)) = h(n, f'(n)) = f'(s(n)),$$

⁷ Per la nozione di sistema di chiusura si veda l'Appendice.

risulta che $s(n) \in X$. Ciò comporta che $X = S$ e che quindi $f = f'$. $\square$

Una importante applicazione del metodo di definizione per ricorsione è il seguente teorema.

Teorema 3.4. La teoria delle terne di Peano è categorica.⁸

Dim. Siano (S, s, z_0) e (S', s', z_0') due terne di Peano. Allora per trovare un omomorfismo dobbiamo trovare una funzione $f : S \rightarrow S'$ tale che,

$$f(z_0) = z_0' \quad ; \quad f(s(x)) = s'(f(x)).$$

Ma tali condizioni costituiscono una definizione per ricorsione e quindi, per quanto abbiamo visto nel Teorema 2.3, esiste una ed una sola funzione f che soddisfa tali condizioni. Tale funzione per il modo in cui è stata definita è un omomorfismo. Non è difficile provare poi che f è un isomorfismo. $\square$

Poiché tutte le terne di Peano sono isomorfe tra loro, non ha importanza quale terna viene considerata. D'ora in poi supponiamo che ne sia stata fissata una, indicheremo con N il suo dominio e con 0 il suo primo elemento. Inoltre chiamiamo *numero naturale* ogni elemento di N .

4. Somma e prodotto in una terna di Peano

E' possibile dare anche una nozione più generale di "definizione per ricorsione" in modo da poter coinvolgere eventualmente insiemi diversi dalla terna di Peano e da poter definire funzioni di più variabili. Ad esempio, consideriamo la funzione potenza n -esima di base b con n numero naturale e b numero reale, consideriamo cioè la funzione $pot(b, n) = b^n$. In questo caso $pot : R \times N \rightarrow R$ è l'unica funzione che soddisfa le condizioni

$$pot(b, 0) = 1 \quad ; \quad pot(b, n+1) = pot(b, n) \cdot b.$$

Tale esempio suggerisce la seguente, più generale, definizione.

⁸ Una teoria si dice *categorica* se tutti i modelli di tale teoria sono isomorfi tra loro. Ad esempio, la teoria dei gruppi non è categorica poiché, come è noto, esistono gruppi che non sono isomorfi tra loro.

Definizione 4.1. Sia (S, s, z_0) una terna di Peano ed A e B insiemi non vuoti. Allora diciamo che la funzione n -aria $f : A \times S \rightarrow B$ è definita per *ricorsione sulla seconda variabile* se esistono due funzioni $g : A \rightarrow B$ ed $h : A \times S^2 \rightarrow B$ tali che:

$$f(a, z_0) = g(a) \quad ; \quad f(a, s(n)) = h(a, y, f(a, n)). \quad (4.1)$$

Tramite tale nozione estesa di ricorsione è possibile definire in una terna di Peano le operazioni di addizione e moltiplicazione.

Definizione 4.2. In una terna di Peano $(N, s, 0)$ chiamiamo *addizione* la funzione $som : N \times N \rightarrow N$ definita per ricorsione tramite le due equazioni:

$$som(x, 0) = x \quad ; \quad som(x, s(y)) = s(som(x, y)).$$

Chiamiamo *moltiplicazione* la funzione $pro : N \times N \rightarrow N$ definita per ricorsione dalle due equazioni

$$pro(x, 0) = 0 \quad ; \quad pro(x, s(y)) = som(pro(x, y), x).$$

In generale la funzione addizione viene indicata con il simbolo $+$ e la funzione di moltiplicazione con un puntino $\cdot$. Inoltre si preferiscono le notazioni *infisse* $x+y$ e $x \cdot y$ al posto delle notazioni *prefisse* $som(x, y)$ e $pro(x, y)$.⁹ Se denotiamo con 1 l'elemento $s(0)$, allora risulta che $som(x, 1) = som(x, s(0)) = s(som(x, 0)) = s(x)$. Pertanto, utilizzando la notazione additiva, possiamo indicare con $x+1$ il successore di x . Per le operazioni ora definite valgono le seguenti proprietà di cui omettiamo la dimostrazione.

Proposizione 4.3. Le operazioni $+$ e $\cdot$ sono associative e commutative ed ammettono come elemento neutro 0 ed 1 , rispettivamente. Inoltre vale la proprietà distributiva del prodotto rispetto la somma.

⁹ Da notare che se provassimo ad estendere le definizioni ora date di addizione e di moltiplicazione ai numeri reali apparirebbero subito delle difficoltà in quanto il campo dei numeri reali non è una terna di Peano. Le difficoltà riguarderebbero sia il processo di calcolo, sia l'unicità della funzione definita. Ad esempio il tentativo di calcolare $som(1, 2.5)$ condurrebbe a calcolare $som(1, 1.5)$, e quindi $som(1, 0.5)$ e quindi $som(1, -0.5)$ e poi $som(1, -1.5)$ e così all'infinito. Per quanto riguarda l'unicità la situazione è la stessa di quella già osservata per il fattoriale.

5. Definire una relazione d'ordine in una terna di Peano.

Vogliamo ora definire una relazione d'ordine in una terna di Peano. Naturalmente le possibili relazioni d'ordine sono infinite ma tra tutte quante vogliamo trovare una relazione $\leq$ tale che:

- il successivo di un numero sia maggiore del numero stesso, cioè tale che $x \leq s(x)$
- $\leq$ sia la più piccola relazione d'ordine con tale proprietà.

In Appendice abbiamo mostrato come si può costruire una tale relazione. Tuttavia anticipiamo quanto fatto in Appendice utilizzando la nozione di composizione iterata di una funzione che si definisce per ricorsione.

Definizione 5.1. Sia $h : A \rightarrow A$ una funzione ed $n \in \mathbb{N}$, allora indichiamo con h^n la funzione definita per ricorsione ponendo, per ogni x in A ,

$$h^0(x) = x ; \quad h^{n+1}(x) = h(h^n(x)).$$

Da notare che, in accordo con la definizione 4.1, la variabile x non ha necessariamente valori in una terna di Peano.

Teorema 5.2. Definiamo la relazione $\leq$ ponendo $x \leq y$ se esiste $n \in \mathbb{N}$ tale che $s^n(x) = y$. Allora $\leq$ è la relazione d'ordine generata dal successore, cioè è la più piccola relazione d'ordine tale che, per ogni $x \in \mathbb{N}$, $x \leq s(x)$.

Dim. Poiché $s^0(x) = x$ abbiamo che $x \leq x$ e quindi $\leq$ verifica la proprietà riflessiva. Inoltre

$$x \leq y, y \leq z \Rightarrow \exists n, m \quad y = s^n(x) \text{ e } z = s^m(y) \Rightarrow z = s^{n+m}(x) \Rightarrow x \leq z.$$

e questo prova che $\leq$ è una relazione transitiva. Per provare che vale la proprietà anti-simmetrica, cominciamo con il provare che, fissato $n \neq 0$, per ogni $x \in \mathbb{N}$ risulta che

$$x \neq s^n(x). \quad (5.1)$$

Ora (5.1) risulta vera per $x = 0$ in quanto 0 non è successore di nessun elemento. Supponiamo che (5.1) sia vera per x , cioè che $x \neq s^n(x)$. Allora, poiché s è una funzione iniettiva, sarà anche $s(x) \neq s(s^n(x)) = s^{n+1}(x)$. Pertanto (5.1) è vera anche per $s(x)$. Per il principio di induzione possiamo concludere che (5.1) è vera per ogni $x \in \mathbb{N}$.

Da (5.1) segue la proprietà antisimmetrica, infatti

$$x \leq y, y \leq x \Rightarrow \exists n, m \quad y = s^n(x) \text{ e } x = s^m(y)$$

$$\Rightarrow x=s^m(s^n(x)) \Rightarrow x=s^{m+n}(x) \Rightarrow n=m=0 \Rightarrow x=y.$$

Abbiamo pertanto provato che $\leq$ è una relazione d'ordine. Per il modo in cui abbiamo definito tale relazione, è immediato che essa contiene l'insieme $\{(x, s(x)) : x \in N\}$ di coppie, cioè che $x \leq s(x)$. Per provare che $\leq$ è la più piccola relazione d'ordine che soddisfa tale proprietà, supponiamo che $\mathcal{R}$ sia una relazione di ordine contenente $\{(x, s(x)) : x \in N\}$. Vogliamo provare che $\mathcal{R}$ contiene $\leq$, cioè che, fissato x ,

$$(x, s^n(x)) \in \mathcal{R} \text{ per ogni } n. \quad (5.2)$$

Procediamo per induzione su n . E' immediato che (5.2) è vera per $n = 0$. Supponiamo che (5.2) sia verificata da n , cioè che $(x, s^n(x)) \in \mathcal{R}$, allora poiché $(s^n(x), s(s^n(x))) \in \mathcal{R}$, per la proprietà transitiva di $\mathcal{R}$ possiamo anche affermare che $(x, s^{n+1}(x)) \in \mathcal{R}$ e quindi (5.2) è vera per $n+1$. Ciò prova che $\mathcal{R}$ contiene $\leq$. $\square$

Chiamiamo *relazione d'ordine naturale* la relazione d'ordine in una terna di Peano generata dalla funzione successore.

Teorema 5.3. La relazione d'ordine naturale definita in una terna di Peano è una relazione totale il cui minimo è 0.

Dim. Per provare che $0 \leq x$ per ogni x , osserviamo che tale disuguaglianza vale per $x = 0$ e che se vale per x vale ovviamente anche per $s(x)$. Per il principio di induzione essa vale per ogni x .

Per provare che $\leq$ è totale consideriamo, dato un elemento x , l'insieme $Conf(x) = \{y \in N : 0 \leq x \leq y \text{ oppure } y \leq x\}$ degli elementi confrontabili con x . E' evidente che $0 \in Conf(x)$. Supponiamo che $y \in Conf(x)$, allora nel caso $x \leq y$ risulta anche che $x \leq y \leq s(y)$ e quindi $s(y) \in Conf(x)$. Consideriamo il caso $y < x$, cioè $x = s^n(y)$ con $n \neq 0$, allora $s^{n-1}(s(y)) = x$ e quindi $s(y) \leq x$. Questo comporta che $s(y) \in Conf(x)$. Per il principio di induzione ciò prova che $Conf(x) = N$ e quindi ogni elemento è confrontabile con x .

Teorema 5.4. La relazione d'ordine naturale definita in una terna di Peano è una relazione di buon ordine¹⁰.

¹⁰ Per la nozione di buon ordine si veda in Appendice.

Dim. Per provare che $\leq$ è un buon ordine cominciamo con il provare che per ogni x non esiste x' tale che $x < x' < s(x)$. A tale scopo osserviamo che in tale caso sarebbe $x' = s^n(x)$ e $s(x) = s^m(x')$ con $n \neq 0$ e $m \neq 0$. Pertanto $s(x) = s^m(s^n(x))$ e ciò è in contrasto con (5.1).

Sia ora X un sottoinsieme non vuoto di N : vogliamo provare che X ammette un minimo. A tale scopo indichiamo con M l'insieme dei minoranti di X . Ovviamente $0 \in M$ e quindi se ogni elemento $m \in M$ avesse il successivo in M , per il principio di induzione avremmo che $M = N$, cioè tutti gli elementi sono minoranti di X . Ma ciò non può accadere in quanto se x è un elemento di X allora $s(x)$ non essendo minore di x non può essere un minorante di X . Ciò comporta che esiste un elemento $m \in M$ tale che $s(m) \notin M$. Affermo che m è il minimo di X . Infatti poiché $s(m)$ non è un minorante di X , esiste $x' \in X$ tale che $s(m)$ non è minore di x' . Poiché $\leq$ è un ordine totale, $x' < s(m)$. Essendo anche $m \leq x'$, ciò implica che $m = x' \in M$ e quindi che m è il minimo di X . $\square$

Possiamo ora dare la seguente fondamentale definizione.

Definizione 5.5. Chiamiamo *sistema di numeri naturali* o *algebra dei numeri naturali* la struttura algebrica ordinata $(N, +, \cdot, 0, \leq)$ che si ottiene definendo in una terna di Peano $(N, s, 0)$ le operazioni di addizione e moltiplicazione e la relazione d'ordine $\leq$.¹¹

Concludiamo con due paradossi che, in contrasto con il teorema 5.3, “dimostrano” come non sia vero che ogni insieme di numeri naturali ammette il minimo.

Paradosso del mucchio di grano.¹² Consideriamo l'insieme X dei numeri naturali che rappresentano il numero di chicchi dei mucchi di grano grandi. E' evidente che X è non vuoto (esiste almeno un mucchio grande) e quindi, se l'insieme dei naturali fosse bene ordinato, X ammetterebbe un minimo m . Ma allora $m-1$, essendo strettamente minore di m non appartiene ad X . Si per-

¹¹ Un modo totalmente diverso di introdurre la struttura algebrica dei numeri naturali è quello che si basa sulla nozione di numero cardinale. Questo punto di vista è esposto nel Capitolo 4.

¹² Abbiamo già esaminato questo paradosso in relazione al principio di induzione. Questa è una riformulazione in termini di buon ordinamento.

viene quindi all'assurdo per cui esiste un mucchio di grano piccolo a cui basta aggiungere un chicco perché diventi grande.

Paradosso di Berry. Chiamo “definizione” di un numero naturale n una frase, scritta nella lingua italiana, del tipo “il numero n tale che ...” dove al posto dei puntini è messa una proprietà verificata da n e solo da n . Ad esempio “il numero n il cui quadrato sia nove” è una definizione di 3. Se ora una definizione la scrivo utilizzando la tastiera di un computer, potrò chiamare “lunghezza della definizione” il numero di volte che ho battuto su tale tastiera per ottenerla (comprendendo spazi vuoti e punti). Ha senso quindi considerare l'insieme A dei numeri naturali che ammettono una definizione di lunghezza minore di 90. Questo insieme è finito poiché è finito l'insieme delle definizioni che contengono meno di 90 battiture.¹³ Supposto per assurdo che N sia ben ordinato, possiamo considerare il minimo m di $B = -A$. Allora m può essere definito come

“il più piccolo tra i numeri che non possono essere definiti da meno di 90 battiture”.

Ma se contiamo i caratteri di tale definizione ci accorgiamo che sono 82. Pertanto abbiamo trovato un modo per definire m con meno di 90 battiture in contrasto con il fatto che $m \notin A$. L'assurdo cui siamo pervenuti prova che l'ordinamento usuale in N non è di buon ordine.

6. Variazioni sul principio di induzione

La definizione della relazione d'ordine permette di riformulare in maniera leggermente più generale il principio di induzione.

Proposizione 6.1. (Principio di induzione). Supponiamo che una proprietà P sia definita in una terna di Peano (S, s, z_0) , che $\underline{m}$ sia un elemento di S e che P verifichi le due affermazioni:

1. P è verificata da $\underline{m}$
2. se P è verificata da x allora è verificata da $s(x)$.

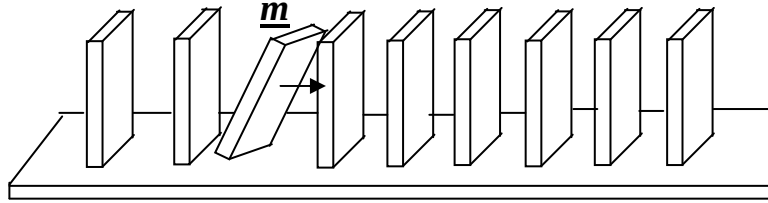
In tale caso è possibile concludere che

P è verificata per ogni $x \geq \underline{m}$.

¹³ Se la tastiera avesse 50 tasti allora l'insieme delle frasi che posso scrivere con 90 battiture avrà cardinalità 50^{90} ed è quindi finito. Sarà allora anche finito l'insieme delle possibili definizioni.

Dim. Sia $X = \{n \in S : n \text{ verifica } P\}$, allora $D = X \cup \{x \in S : x < \underline{m}\}$ contiene z_0 ed è tale che $s(D) \subseteq D$. Pertanto $D = S$ e quindi $X \supseteq \{x \in S : x \geq \underline{m}\}$. $\square$

Questa forma del principio di induzione può essere visualizzata dalla seguente figura il cui significato è ovvio:



Problema. Supponiamo che in un gioco come il poker abbiamo a disposizione solo gettoni che valgono 3 oppure 5 euro. Dimostrare che un giocatore può mettere nel piatto qualsiasi puntata maggiore di 7 euro.

Altre variazioni del principio di induzione si possono ottenere introducendo le terne di Peano ed il principio di induzione in termini più algebrici. Ricordiamo che viene chiamata *parte stabile* di una struttura algebrica A ogni sottoinsieme di A che contenga gli elementi designati e che sia chiuso rispetto alle operazioni della struttura. L'insieme delle parti stabili è un sistema di chiusura, cioè l'intersezione di una famiglia di parti stabili di una struttura è ancora una parte stabile. Ciò permette, dato un è un sottoinsieme X di A di definire la *parte stabile generata da X* come l'intersezione di tutte le parti stabili contenenti X . Indichiamo con $\langle X \rangle$ tale parte. Se $\langle X \rangle = A$, allora si dice che X è un sistema di generatori di A . Ad esempio in un gruppo $(G, \cdot, ^{-1}, 1)$ una parte stabile è un sottoinsieme G' di G tale che

- G contiene 1
- il prodotto di due elementi di G' appartiene ancora a G' ,
- l'inverso di un elemento in G' appartiene ancora a G' .

In definitiva le parti stabili di G coincidono con i sottogruppi di G . Se G può essere generato da un solo elemento, allora si dice che G è ciclico.

Proposizione 6.2. Una struttura algebrica (S, s, z_0) soddisfacente P_1 e P_2 è una terna di Peano se e solo se ammette z_0 come generatore, cioè se S è la più piccola parte stabile contenente z_0 .

Dim. Ovvio perché il principio di induzione afferma proprio che ogni parte X che sia stabile rispetto all'operazione s e che contenga z_0 coincide con S . $\square$

Questo modo di vedere le terne di Peano permette di provare in modo più rigoroso l'esistenza di terne di Peano. Ricordiamo la definizione di insieme infinito.

Definizione 6.3. Chiamiamo *infinito* un insieme T che sia equipotente ad una sua parte propria, cioè tale che esista una funzione iniettiva $f: T \rightarrow T$ che non sia suriettiva, cioè $f(T) \neq T$.

L'esistenza di una terna di Peano equivale ad accettare l'esistenza di un insieme infinito.

Teorema 6.4. Esiste una terna di Peano se e solo se esiste un insieme infinito.

Dim. Supponiamo che esista una terna di Peano (S, s, z_0) , allora essendo la funzione successore iniettiva e poiché $z_0 \notin s(S)$, S è un insieme infinito. Viceversa sia T un insieme infinito e sia $f: T \rightarrow T$ una funzione iniettiva tale che $f(T) \neq T$. Allora esiste un elemento $z_0 \notin f(T)$ e possiamo prendere in considerazione la struttura algebrica (T, f, z_0) . Sia $S = \langle z_0 \rangle$ la parte stabile generata da z_0 , vogliamo provare che la sottostruttura (S, f, z_0) è una terna di Peano.¹⁴ Infatti, gli assiomi P_1 e P_2 sono evidenti. Per provare il principio di induzione basta osservare che S , per costruzione, è la più piccola parte stabile contenente z_0 . $\square$

La Proposizione 6.2 suggerisce la seguente estensione del principio di induzione la cui dimostrazione è evidente.

Teorema 6.5. Sia $(S, h_1, \dots, h_t, z_0, \dots, z_m)$ una struttura algebrica avente $z_0, \dots, z_m$ come sistema di generatori. Supponiamo inoltre che P sia una proprietà tale che:

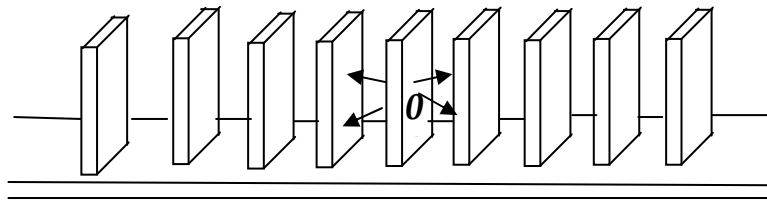
¹⁴ Per la nozione di sottostruttura generata si veda l'Appendice.

- P è verificata da $z_0, \dots, z_m$
 - per ogni operazione n -aria h_i se P è verificata da $x_1, \dots, x_n$ allora è verificata da $h_i(x_1, \dots, x_n)$.
- Allora P è verificata per ogni $x \in S$.

Esempio: “Principio di induzione” per $\mathbb{Z}$. Come già osservato, la struttura $(\mathbb{Z}, s, 0)$, dove $s(x) = x+1$, non è una terna di Peano e quindi in $\mathbb{Z}$ non sarebbe possibile fare dimostrazioni per induzione. Tuttavia se vogliamo provare che una proposizione vale per ogni elemento di $\mathbb{Z}$ possiamo riferirci alla struttura algebrica $(\mathbb{Z}, s_+, s_-, 0)$ dove $s_+(x) = x+1$ e $s_-(x) = x-1$. Infatti la sottostruttura di $(\mathbb{Z}, s_+, s_-, 0)$ generata da 0 coincide con $(\mathbb{Z}, s_+, s_-, 0)$, cioè 0 è un generatore di $(\mathbb{Z}, s_+, s_-, 0)$. Se volessimo formulare un principio di induzione per $\mathbb{Z}$ dovremmo quindi dire che, per ogni proprietà P definita in $\mathbb{Z}$,

- se P è verificata da 0
- e se P è verificata da x comporta che P sia verificata da $x+1$ e da $x-1$,
- allora P è verificata per ogni $x \in \mathbb{Z}$.

In tale caso l'applicazione del principio di induzione può essere visualizzata al modo seguente. Consideriamo la seguente figura in cui i pezzi del gioco domino sono poggiati su di un tavolo infinito sia a sinistra che a destra:



Supponiamo che tali pezzi siano esplosivi e che:

- il pezzo **0** esplode
 - se un pezzo x esplode allora fa esplodere i due pezzi vicini, cioè sia il pezzo successivo $s_+(x)$ che quello precedente $s_-(x)$
- allora è evidente che tutti i pezzi esplodono.

Esempio: il principio di induzione in $\mathbb{Z}/m$. Abbiamo già osservato che la struttura $(\mathbb{Z}/m, s, [0])$, dove s è definita ponendo $s([n]) = [n+1]$ non è una terna di Peano. Tuttavia è vero che $[0]$ è un generatore di tale struttura poiché ogni intero modulo m si può

ottenere a partire da $[0]$ applicando un certo numero di volte s . Pertanto pur non essendo $(\mathbb{Z}/m, s, [0])$ una terna di Peano continua a valere il principio di induzione.

7. L'anello degli interi relativi

I numeri naturali sono uno strumento per misurare la grandezza di insiemi finiti. Ad esempio la terna di Peano delle possibili tacche su di un pezzo di legno può avere come scopo il contare il numero di pecore o il numero dei giorni passati od altro. Tuttavia esistono tipi di attività in cui i numeri naturali si mostrano inadeguati. Supponiamo ad esempio di dovere distinguere in una contabilità i soldi che devono essere dati dai soldi che si devono ricevere da alcuni

clienti. Se la contabilità è tenuta su due colonne avremo una situazione del tipo indicato nella tabella.

Cliente	Avere	Dare
Carlo	7	5
Luigi	3	6
Maria	10	8
...	...	...

Ne segue che l'informazione relativa a Carlo è rappresentata dalla coppia $(7,5)$, l'informazione relativa a Luigi è rappresentata dalla coppia $(3,6)$, quella relativa a Maria da $(10,8)$. Ciò suggerisce l'introduzione di un nuovo tipo di numero costituito da due parti (quindi una coppia) che hanno significato diverso. Naturalmente oltre all'interpretazione di una coppia in termini di debiti e crediti, sono possibili diverse altre interpretazioni. Ad esempio possiamo interpretare una coppia (n,m)

- come l'operazione "aggiungere n e togliere m ".
- come "fare n passi avanti ed m passi indietro"
- come "applicare una forza di grandezza n in una direzione ed una forza di grandezza m nella direzione opposta".

In definitiva partiamo dall'insieme dei numeri naturali N (compreso lo zero) e consideriamo l'insieme $N \times N$. In tale insieme definiamo una operazione di addizione ponendo

$$(n,m) + (a,b) = (n+a, m+b).$$

Il motivo per cui la somma viene definita in questo modo è ovvio. Se si fanno n passi avanti ed m indietro e poi si fanno a passi avanti e b indietro, allora globalmente si sono fatti $n+a$ passi avanti ed $m+b$ indietro.

La definizione di somma, che è associativa, permette di definire il multiplo n -esimo di un elemento (a,b) come l'elemento che si ottiene sommando n volte a se stessa la coppia (a,b) . Ciò permette di scomporre ogni coppia (m,n) al modo seguente

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

Se si indica con $+1$ la coppia $(1,0)$ e con -1 la coppia $(0,1)$ possiamo scrivere tale scomposizione al modo seguente:

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

L'operazione di moltiplicazione viene invece definita al modo seguente:

$$(n,m) \cdot (a,b) = (na+mb, nb+ma). \quad (7.1)$$

Giustificare una tale definizione è alquanto difficile perché, ad esempio, non ha molto senso moltiplicare passi avanti con passi indietro (un problema analogo si presenta quando si deve definire la moltiplicazione tra numeri complessi). Tuttavia anche se non è chiara la giustificazione "semantica" della moltiplicazione, ne esiste una "sintattica" nel senso che (7.1) è la sola possibile definizione se vogliamo ottenere che il prodotto sia una operazione distributiva rispetto alla somma e che valga qualche ovvia proprietà.

Proposizione 7.1. Supponiamo di volere definire una operazione $\cdot$ in $N \times N$ in modo che:

- i) valga la proprietà distributiva di $\cdot$ rispetto a $+$
- ii) $+1$ sia elemento neutro,
- iii) $(-1)^2 = 1$.

Allora (7.1) è l'unica possibile definizione.

Dim. Sia $*$ una operazione binaria che verifica i), ii) e iii). Allora

$$\begin{aligned} (n,m) * (a,b) &= [(n,0) + (0,m)] * (a,b) = (n,0) * (a,b) + (0,m) * (a,b) \\ &= n \cdot [(1,0) * (a,b)] + m \cdot [(0,1) * (a,b)] \\ &= n \cdot [(a,b)] + m \cdot (b,a) \\ &= (na+mb, nb+ma) = (n,m) \cdot (a,b) \end{aligned}$$

e quindi che $*$ e $\cdot$ coincidono.

Viceversa, supponiamo che $\cdot$ sia definita tramite l'equazione (7.1), allora è facile verificare che $\cdot$ soddisfa i), ii) e iii). $\square$

Proposizione 7.2. Data la struttura $(N \times N, +, \cdot, 0, 1)$ le operazioni $+$ e $\cdot$ sono associative e commutative ed ammettono $0 = (0,0)$ e $1 = (1,0)$ come elemento neutro, rispettivamente. Inoltre vale la

proprietà distributiva. Tuttavia $(N \times N, +, \mathbf{0})$ non è un gruppo e quindi $(N \times N, +, \cdot, \mathbf{0}, \mathbf{1})$ non è un anello.

Dim. Ci limitiamo ad osservare che, dato un elemento (m, n) diverso da $(0, 0)$, qualunque sia (a, b) risulta che

$$(m, n) + (a, b) = (m+a, n+b) \neq (0, 0).$$

Pertanto in $(N \times N, +, \mathbf{0})$ non esiste l'opposto di (m, n) .

La struttura $(N \times N, +, \cdot, \mathbf{0}, \mathbf{1})$ non è adeguata a rappresentare le situazioni da cui siamo partiti all'inizio del paragrafo per il fatto che elementi diversi di $N \times N$ possono rappresentare la stessa situazione. Ad esempio, se si interpretano le coppie in termini di debiti e di crediti, allora è naturale considerare (n, m) equivalente a (n', m') se avere n e dare m risulta equivalente ad avere n' e dare m' . Ad esempio se si guarda la tabella, è evidente che Carlo e Maria sono due clienti con la stessa situazione finanziaria, cioè che la coppia $(7, 5)$ è equivalente alla coppia $(10, 8)$. Ancora, è naturale considerare (n, m) equivalente a (n', m') se fare n passi avanti ed m indietro produce lo stesso risultato di fare n' passi avanti ed m' indietro. Potremmo allora dire che (n, m) è equivalente a (n', m') se nel caso $n \geq m$ risulta che $n' \geq m'$ e $n - m = n' - m'$, mentre nel caso $n < m$ risulta che $n' < m'$ e $m - n = m' - n'$.

Proposizione 7.3. Definiamo la relazione $\equiv$ in $N \times N$ ponendo,

$$(n, m) \equiv (n', m') \Leftrightarrow n + m' = m + n'. \quad (7.2)$$

Allora tale relazione è una congruenza della struttura algebrica $(N \times N, +, \cdot, (0, 0), (1, 0))$. Il relativo quoziente è un anello unitario.

Dim. Proviamo che

$$(n, m) \equiv (n', m'), (a, b) \equiv (a', b') \Rightarrow (n+a, m+b) \equiv (n'+a', m'+b').$$

Infatti per ipotesi $n+m' = m+n'$ e $a+b' = b+a'$, da cui, sommando termine a termine, $n+m'+a+b' = m+n'+b+a'$ che equivale a $(n+a, m+b) \equiv (n'+a', m'+b')$. La dimostrazione della compatibilità rispetto al prodotto e del fatto che il quoziente sia un anello viene lasciata come esercizio al lettore. $\square$

Definizione 7.4. Chiamiamo *anello degli interi relativi* la struttura quoziente di $(N \times N, +, \cdot, \mathbf{0}, \mathbf{1})$ modulo $\equiv$. Indichiamo con $(\mathbb{Z}, +, \cdot, 0, 1)$ tale struttura.

Precisamente Z è definita dalle equazioni

$$\begin{aligned} [(n,m)] &= \{(n',m') \mid (n',m') \equiv (n,m)\} \\ Z &= \{[(n,m)] \mid (n,m) \in \mathbb{N} \times \mathbb{N}\} \\ [(n,m)] + [(a,b)] &= [(n+a, m+b)] \\ [(n,m)] \cdot [(a,b)] &= [(na+mb, nb+ma)] \\ 0 &= [\mathbf{0}] \\ 1 &= [\mathbf{1}]. \end{aligned}$$

Esercizio. Provare che $[(1,0)]$ è l'elemento neutro rispetto al prodotto in Z .

Esercizio. Dire perché è sbagliato definire l'operazione $\otimes$ ponendo $[(n,m)] \otimes [(a,b)] = [(na, mb)]$.

8. Il campo dei razionali.

Il passaggio dall'anello Z degli interi relativi al campo Q dei numeri razionali si ottiene in modo analogo a quello del passaggio dai naturali agli interi relativi. In questo caso consideriamo l'insieme

$$Z \times (Z - \{0\}) = \{(p,q) \mid p \in Z, q \in Z, q \neq 0\}$$

e l'interpretazione che ora diamo ad una coppia (p,q) è "moltiplicare per p e dividere per q ". In tale insieme di coppie introduciamo due operazioni tramite le eguaglianze

$$(p,q) + (a,b) = (pb+qa, qb) ; \quad (p,q) \cdot (a,b) = (pa, qb) \quad (8.1)$$

ottenendo la seguente struttura algebrica

$$(Z \times (Z - \{0\}), +, \cdot, (0,1), (1,1)).$$

Proposizione 8.1. Nella struttura $(Z \times (Z - \{0\}), +, \cdot, (0,1), (1,1))$ le operazioni sono commutative ed associative, $(0,1)$ è elemento neutro rispetto a $+$, $(1,1)$ è elemento neutro rispetto a $\cdot$. Tuttavia tale struttura non è un campo.

Dim. Poichè $(p,q) + (0,1) = (p1+q0, q1) = (p,q)$, la coppia $(0,1)$ è elemento neutro rispetto la somma. In modo simile si provano le altre proprietà. Per provare che la struttura non è un campo osserviamo che se una coppia (p,q) ammettesse inverso allora esisterebbero due interi x ed y in Z tali che $(p,q) \cdot (x,y) = (1,1)$, si avrebbe pertanto che $px = 1$ e $qy = 1$, e quindi p e q sarebbero invertibili in Z . Poiché gli unici elementi invertibili di Z sono 1 ed il suo opposto -1, ne segue che gli unici elementi invertibili della

struttura $(Z \times (Z - \{0\}), +, \cdot, (0,1), (1,1))$ sono le coppie $(1,1)$ e $(-1,-1)$, $(1,-1)$ e $(-1,1)$ che ammettono come inverso se stesse. $\square$

Possiamo ottenere un campo da $(Z \times (Z - \{0\}), +, \cdot, (0,1), (1,1))$ introducendo una opportuna congruenza e passando a quoziente. Ancora una volta possiamo ritenere equivalenti due coppie se “producono lo stesso effetto”, allora, ad esempio, dobbiamo identificare la coppia $(3,4)$ con la coppia $(6,8)$. Per convincersi di questo fatto basta ricorrere alle solite torte che vengono proposte ai bambini a cui si insegnano le frazioni ed accorgersi che tre quarti di una torta coincidono con i sei ottavi di una torta. Ciò conduce a definire la seguente relazione.

Proposizione 8.2. Sia $\equiv$ la relazione in $Z \times (Z - \{0\})$ definita ponendo,

$$(p,q) \equiv (p',q') \Leftrightarrow p \cdot q' = q \cdot p'. \quad (8.2)$$

Tale relazione è una equivalenza compatibile con le operazioni di somma e prodotto date in (8.1) ed è pertanto una congruenza.

Dim. Per verificare la compatibilità con il prodotto osserviamo che

$$\begin{aligned} (n,m) \equiv (n',m'), (p,q) \equiv (p',q') &\Rightarrow n \cdot m' = m \cdot n', p \cdot q' = q \cdot p' \\ &\Rightarrow n \cdot m' \cdot p \cdot q' = m \cdot n' \cdot q \cdot p' \Leftrightarrow (n \cdot p, m \cdot q) \equiv (n' \cdot p', m' \cdot q') \\ &\Leftrightarrow (n,m) \cdot (p,q) \equiv (n',m') \cdot (p',q'). \end{aligned}$$

Similmente si prova la compatibilità rispetto la somma. $\square$

Proposizione 8.3. La struttura $(Q, +, \cdot, 0, 1)$ ottenuta come quoziente modulo $\equiv$ di $(Z \times (Z - \{0\}), +, \cdot, (0,1), (1,1))$ è un campo che chiamiamo *campo dei numeri razionali*.

Dim. Per provare che $(Q, +, \cdot, 0, 1)$ è un campo ricordiamo che tale struttura è definita dalle equazioni

$$\begin{aligned} [(p,q)] &= \{(p',q') \mid (p',q') \equiv (p,q)\} \\ Q &= \{[(p,q)] \mid (p,q) \in N \times N, q \neq 0\} \\ [(p,q)] + [(a,b)] &= [(p \cdot b + q \cdot a, q \cdot b)] \\ [(p,q)] \cdot [(a,b)] &= [(p \cdot a, q \cdot b)]. \\ 0 &= [(0,1)] \\ 1 &= [(1,1)]. \end{aligned}$$

E' immediato verificare che $[(0,1)]$ è elemento neutro rispetto all'addizione e che $[(1,1)]$ è elemento neutro rispetto alla multi-

plicazione. Allora per provare che $[(p,q)]$ ammette opposto, osserviamo che $[(p,q)] + [(-p,q)] = [(p \cdot q - p \cdot q, q \cdot q)] = [(0, q \cdot q)]$. D'altra parte, poiché $(0,1) \equiv (0, q \cdot q)$, la classe $[(0, q \cdot q)]$ coincide con la classe $[(0,1)] = 1$. In conclusione l'opposto di $[(p,q)]$ è $[(-p,q)]$. Ad esempio il numero $[(3,4)]$ ammette come inverso $[(-3,4)]$ poiché $[(3,4)] + [(-3,4)] = [(12 + -12, 16)] = [(0,16)]$ e, poiché $(0,16) \equiv (0,1)$, risulta che $[(0,16)] = [(0,1)]$.

Sia $[(p,q)]$ non nullo e cioè tale che $p \neq 0$, allora $[(q,p)]$ è l'inverso di $[(p,q)]$. Infatti, essendo $(p \cdot q, q \cdot p)$ equivalente a $(1,1)$, risulta che

$$[(p,q)] \cdot [(q,p)] = [(p \cdot q, q \cdot p)] = [(1,1)].$$

Ad esempio $[(3,4)]$ ammette come inverso $[(4,3)]$ poiché $[(3,4)] \cdot [(4,3)] = [(12,12)]$ e, poiché $(12,12) \equiv (1,1)$. Le rimanenti proprietà di campo sono semplici da dimostrare. $\square$

Problema. Poiché siamo liberi di definire le operazioni in un insieme per inventare nuove strutture algebriche, definiamo nell'insieme dei numeri razionali l'operazione $\oplus$ ponendo

$$n/m \oplus p/q = (n^2 + p^2) / (m^2 + q^2).$$

1. Dire se l'operazione definita in questo modo è commutativa.
2. Dire perché è sbagliato chiedersi se l'operazione definita in questo modo è commutativa.

Concludiamo dicendo che nel campo dei razionali è definita una relazione d'ordine che lo rende un campo ordinato. Omettiamo la dimostrazione che si riduce ad una semplice verifica.

Proposizione 8.4. Definiamo nel campo dei razionali una relazione $\leq$ definita ponendo

$$[(p,q)] \leq [(n,m)] \Leftrightarrow p \cdot m \leq q \cdot n.$$

Allora il campo dei razionali con tale relazione diviene un campo ordinato.

9. I numeri reali tramite le sezioni

Più delicato, da un punto di vista filosofico, è il passaggio dai razionali ai reali. Esponiamo, ad esempio, il metodo delle sezioni di Dedekind che è quello più utilizzato per la costruzione del campo dei reali anche se a mio parere è alquanto ferraginoso. Torniamo alla teoria delle grandezze omogenee che abbiamo esposto nel primo capitolo e supponiamo che in una classe $(G, =$

, <, +) di grandezze omogenee sia stata fissata una unità di misura $u \in G$. Allora, come abbiamo già osservato nel primo capitolo, se g è una grandezza da misurare un primo tentativo di misurazione consisterà nel prendere multipli successivi di u fino a raggiungere g . Se si trova un intero p tale che $p \cdot u = g$ allora è possibile concludere che la misura di g rispetto ad u è p . Altrimenti si considera un naturale p tale che $p \cdot u < g < (p+1) \cdot u$ ed in tale caso si dice che p è una misura per difetto e $p+1$ una per eccesso di g . Una misurazione più precisa si può comunque avere dividendo u in q parti uguali (in generale in dieci parti) ed assumendo come sotto-unità di misura $u' = u/q$. Ora potrebbe capitare che per un opportuno p risulti che $p \cdot u' = (p/q) \cdot u = g$. In tale caso si concluderebbe che la misura cercata è p/q . Se invece ciò non accade allora potremmo lo stesso trovare p tale che $p \cdot u' < g < (p+1) \cdot u'$ e quindi $(p/q) \cdot u < g < (p+1)/q \cdot u$ e concludere che p/q è una misura per difetto e $(p+1)/q$ una misura per eccesso di g . Una misura più precisa si può avere dividendo la nuova unità di misura u' in un numero abbastanza alto di parti uguali. Nel caso in cui u e g siano incommensurabili, cioè che non esista un razionale p/q tale che $g = (p/q) \cdot u$ (come nel caso della diagonale e del lato del quadrato) tale processo di approssimazione non finisce mai. Allora in tale caso possiamo comunque considerare l'insieme A dei razionali positivi che misurano per difetto g e l'insieme B dei razionali positivi che misurano per eccesso g

$$A = \{p/q \in Q^+ \mid (p/q) \cdot u < g\}, \quad B = \{p/q \in Q^+ \mid (p/q) \cdot u > g\}.$$

E' facile verificare che:

- a) $A \cap B = \emptyset$
- b) $A \cup B = Q^+$
- c) $x \in A, y \leq x \Rightarrow y \in A$; $x \in B, y \geq x \Rightarrow y \in B$.
- d) A è privo di massimo, B è privo di minimo.

Il metodo delle sezioni in un certo senso chiama “numero irrazionale positivo” una coppia di sottoinsiemi di Q^+ di questo tipo.¹⁵

¹⁵Ecco quanto afferma Dedekind a tale proposito:

Ora, in ogni caso in cui c'è una sezione (A, B) che non è prodotta da un numero razionale, allora noi creiamo un nuovo numero irrazionale che riteniamo completamente definito da questa sezione; diremo che questo numero corrisponde a questa sezione oppure che produce questa sezione.

Definizione 9.1. Una *numero irrazionale positivo* è una coppia (A, B) di sottoinsiemi di Q^+ verificanti le condizioni *a)*, *b)*, *c)* e *d)*.

Ora il nostro scopo è immergere sia i razionali che gli irrazionali in un unico ambiente in modo da potere parlare in generale di numero reale. A tale scopo possiamo identificare un numero razionale r con la coppia (A_r, B_r) di insiemi di razionali dove

$$A_r = \{x \in Q^+ \mid x < r\} \quad ; \quad B_r = \{x \in Q^+ \mid x > r\}.$$

In questo caso sono verificate le proprietà *a)*, *c)* e *d)* mentre al posto della proprietà *b)* risulta che $A \cup B = Q^+ - \{r\}$. In altri termini se u e g sono commensurabili, ad esempio $u = r \cdot g$, per questione di uniformità di notazione indicheremo con (A_r, B_r) la misura di g rispetto ad u . In ogni caso viene individuata una coppia (A, B) di sottoinsiemi di Q^+ in cui, ripetiamo, la prima componente è vista come l'insieme delle misure per difetto e la seconda componente come l'insieme delle misure per eccesso di g .

Definizione 9.2. Chiamiamo *sezione positiva* o *numero reale positivo* una coppia (A, B) di insiemi di razionali tali che:

- a)* $A \cap B = \emptyset$
- b)* $A \cup B = Q^+$ oppure esiste $r \in Q^+$ tale che $A \cup B = Q^+ - \{r\}$
- c)* $x \in A, y \leq x \Rightarrow y \in A \quad ; \quad x \in B, y \geq x \Rightarrow y \in B.$
- d)* A è privo di massimo, B è privo di minimo.

Diremo che (A, B) è *irrazionale* se è verificata la prima delle condizioni in *b)*, che (A, B) è *razionale* se è verificata la seconda delle condizioni.

Nel seguito indicheremo

- con 1 la sezione $(\{x \in Q^+ \mid x < 1\}, \{x \in Q^+ \mid x > 1\})$
- con 0 la sezione $(\{x \in Q^+ \mid x < 0\}, \{x \in Q^+ \mid x > 0\})$
- ... e così via.

In altri termini identificheremo un razionale con la sezione da esso determinata. Un esempio di sezione che non è un razionale è il seguente

$$A = \{x \in Q^+ \mid x^2 < 2\} \quad ; \quad B = \{x \in Q^+ \mid x^2 > 2\}$$

che, in un certo senso, rappresenta il numero $\sqrt{2}$.

Per definire le operazioni aritmetiche tra numeri reali positivi, dati due insiemi X ed Y di numeri reali, poniamo

$$X+Y = \{x+y \mid x \in X, y \in Y\} ; X \cdot Y = \{x \cdot y \mid x \in X, y \in Y\}.$$

Definizione 9.3. Chiamiamo *struttura algebrica dei reali positivi* la struttura $\mathbf{R}^+ = (R^+, +, \cdot)$ dove

- R^+ è l'insieme delle sezioni positive
- l'operazione $+$ è definita ponendo:

$$(A, B) + (A', B') = (A+A', B+B')$$
- l'operazione $\cdot$ è definita ponendo:

$$(A, B) \cdot (A', B') = (A \cdot A', B \cdot B').$$

Inoltre si definisce un ordinamento $\leq$ ponendo

$$(A, B) \leq (A', B') \Leftrightarrow a \leq b' \text{ per ogni } a \in A \text{ e } b' \in B'.$$

Proposizione 9.4. La somma ed il prodotto di due sezioni è ancora una sezione.

Dim. La dimostrazione è lunga e noiosa e non viene fatta. Proviamo solo, per fare un esempio, che $(A \cdot A', B \cdot B')$ verifica la condizione a). Infatti osserviamo che se $a \in A$ e $b \in B$, $a' \in A'$, $b' \in B'$, allora essendo il prodotto strettamente crescente (sui razionali positivi) e risultando che $a < b$ e $a' < b'$, possiamo asserire che $a \cdot a' < b \cdot b'$. Pertanto $(A \cdot A') \cap (B \cdot B') = \emptyset$. la condizione b) supponiamo per assurdo che esista $x \in (A \cdot A') \cap (B \cdot B')$, cioè che esistano $a, a' \in A$ e $b, b' \in B$ tali che $x = a \cdot a' = b \cdot b'$. Poichè quindi $a = (a')^{-1} \cdot b \cdot b'$.

a) $A \cap B = \emptyset$

b) $A \cup B = Q^+$ oppure esiste $r \in Q^+$ tale che $A \cup B = Q^+ - \{r\}$

c) $x \in A, y \leq x \Rightarrow y \in A$; $x \in B, y \geq x \Rightarrow y \in B$.

d) A è privo di massimo, B è privo di minimo.

In tale modo si definisce la struttura algebrica dei reali positivi. Successivamente si procede alla simmetrizzazione di tale struttura con un metodo analogo a quello che ha permesso di costruire l'anello degli interi Z a partire dai numeri naturali. Tuttavia non mi soffermo nei particolari di tale procedimento poiché mi sembra preferibile definire i reali tramite il metodo esposto nel prossimo paragrafo.¹⁶

¹⁶ Usualmente nei libri di testo la nozione di sezione viene definita a partire dall'intero insieme Q dei razionali. Ciò permette di evitare il processo di simmetrizzazione. Tuttavia in tale caso la definizione del prodotto diviene poco naturale e noiosa per il fatto che il prodotto di due razionali negativi è un razionale positivo.

10. I numeri reali tramite le successioni di Cauchy

Il metodo delle sezioni per definire i reali anche se è perfetto da un punto di vista formale, non corrisponde molto all'esperienza di chi si trova effettivamente a manipolare tali numeri. Infatti quando si utilizza un numero reale o lo rappresenta come espansione decimale infinita (quindi come serie di potenze) oppure, più in generale, tramite una successione di razionali il cui limite è il numero reale in questione. Una definizione del campo dei numeri reali che è molto più vicina a questo modo di procedere si ottiene al modo seguente.¹⁷ Indichiamo con $(Q^N, +, \cdot, \underline{0}, \underline{1})$ la potenza diretta di Q con insieme di indici N . Tale struttura è definita assumendo che:

- il dominio è l'insieme Q^N delle successioni di numeri razionali,
- l'addizione $+$ è definita ponendo

$$(a_n)_{n \in N} + (b_n)_{n \in N} = (a_n + b_n)_{n \in N}$$
- la moltiplicazione è definita ponendo;

$$(a_n)_{n \in N} \cdot (b_n)_{n \in N} = (a_n \cdot b_n)_{n \in N}$$
- $\underline{0}$ denota la successione $(z_n)_{n \in N}$ con z_n costantemente uguale a 0
- $\underline{1}$ denota la successione $(u_n)_{n \in N}$ con u_n costantemente uguale a 1.

Proposizione 10.1. La struttura $(Q^N, +, \cdot, \underline{0}, \underline{1})$ è un anello unitario che non è un campo.

Dim. Per provare la proprietà commutativa della somma, che si esprime con una equazione del tipo $x+y = y+x$, osserviamo che

$$(a_n)_{n \in N} + (b_n)_{n \in N} = (a_n + b_n)_{n \in N} = (b_n + a_n)_{n \in N} = (b_n)_{n \in N} + (a_n)_{n \in N}.$$

Per provare che $\underline{0} = (z_n)_{n \in N}$ è l'elemento neutro rispetto all'addizione, osserviamo che

$$(a_n)_{n \in N} + (z_n)_{n \in N} = (a_n + z_n)_{n \in N} = (a_n)_{n \in N}.$$

Per provare che $\underline{1} = (u_n)_{n \in N}$ è l'elemento neutro rispetto alla moltiplicazione, osserviamo che

$$(a_n)_{n \in N} \cdot (u_n)_{n \in N} = (a_n \cdot u_n)_{n \in N} = (a_n)_{n \in N}.$$

¹⁷ Tale definizione è dovuta a Cantor che, non dimentichiamolo, ha sviluppato la teoria degli insiemi a partire da ricerche legate all'analisi matematica. La tecnica è la stessa con cui si costruisce il completamento di uno spazio metrico.

Gli altri assiomi di teoria degli anelli si dimostrano in modo altrettanto banale. Per dimostrare che Q^N non è un campo consideriamo $(a_n)_{n \in N}$ una successione $(a_n)_{n \in N}$ diversa dalla successione nulla ma che abbia un elemento uguale a zero, ad esempio $a_k = 0$. Allora tale successione non può ammettere inverso in quanto comunque si scelga una successione $(b_n)_{n \in N}$ risulterà che $a_k \cdot b_k = 0$ e ciò comporta che $(a_n)_{n \in N} \cdot (b_n)_{n \in N}$ è necessariamente diverso da $(u_n)_{n \in N}$.¹⁸ $\square$

Per potere ottenere un campo, come vedremo, dobbiamo considerare prima una opportuna sottostruttura di $(Q^N, +, \cdot, \underline{0}, \underline{1})$ e poi un quoziente.

Definizione 10.2. Chiamiamo *successione di Cauchy* un elemento $(r_n)_{n \in N}$ in Q^N tale che:

$$\forall \varepsilon > 0 \exists m \forall p \geq m \forall q \geq m |r_p - r_q| \leq \varepsilon.$$

Indichiamo con Ch l'insieme delle successioni di Cauchy.

Tutte le successioni costantemente uguali ad un dato numero razionale r sono esempi di successioni di Cauchy. In particolare sono successioni di Cauchy la successione costantemente uguale a zero e quella costantemente uguale ad 1, cioè $\underline{0}$ ed $\underline{1}$.

Proposizione 10.3. L'insieme Ch delle successioni di Cauchy è un sottoanello di $(Q^N, +, \cdot, \underline{0}, \underline{1})$. Tale anello non è un campo.

Dim. Poiché si dimostra che la somma ed il prodotto di due successioni di Cauchy è una successione di Cauchy, Ch è una parte stabile di Q^N e quindi è ancora un anello unitario. Se considera-

¹⁸Un modo più generale per provare che Q^N è un anello deriva da considerazioni di algebra universale. Infatti gli assiomi che caratterizzano l'essere un anello unitario sono tutti espressi tramite equazioni ed in algebra universale si prova che se una equazione vale per una famiglia di strutture allora vale anche per il prodotto diretto di questa famiglia. Pertanto ogni potenza diretta di un anello unitario è ancora un anello unitario. In particolare la potenza diretta Q^N è un anello unitario. D'altra parte l'aver provato che Q^N non è un campo mostra che la proprietà di essere campo, che usualmente si rappresenta con l'asserzione $\forall x(x \neq 0 \rightarrow \exists y(x \cdot y = 1))$ non può essere espressa da una equazione.

mo la successione $(a_n)_{n \in \mathbb{N}}$ definita ponendo $a_n = 1/n$ se n è pari e $a_n = 0$ se n è dispari, allora abbiamo un esempio di successione di Cauchy che non ammette inverso pur non essendo la successione nulla. Infatti comunque si consideri una successione $(b_n)_{n \in \mathbb{N}}$ è evidente che $(a_n \cdot b_n)_{n \in \mathbb{N}}$ è diversa dalla successione costantemente uguale ad 1.¹⁹

Definizione 10.4. Diciamo che due successioni di Cauchy $(a_n)_{n \in \mathbb{N}}$ e $(b_n)_{n \in \mathbb{N}}$ sono *equiconvergenti* e poniamo $(a_n)_{n \in \mathbb{N}} \equiv (b_n)_{n \in \mathbb{N}}$ se $\lim_{n \rightarrow \infty} |a_n - b_n| = 0$.

Allora due successioni $(a_n)_{n \in \mathbb{N}}$ e $(b_n)_{n \in \mathbb{N}}$ sono equiconvergenti se

$$\forall \varepsilon > 0 \exists m \forall p \geq m |a_p - b_p| \leq \varepsilon.$$

In particolare sono equiconvergenti a $\underline{0}$ tutte e sole le successioni convergenti a zero, sono equiconvergenti a $\underline{1}$ tutte e sole le successioni convergenti ad 1.

Proposizione 10.5. La relazione di equiconvergenza $\equiv$ è una congruenza nell'anello $(Ch, +, \cdot, \underline{0}, \underline{1})$. Il corrispondente quoziente $(Ch/\equiv, +, \cdot, [\underline{0}], [\underline{1}])$ è un campo.²⁰

Dim. E' evidente che $\equiv$ è una relazione di equivalenza. Per provare che è compatibile con il prodotto, supponiamo che $(a_n)_{n \in \mathbb{N}}$ sia equiconvergente a $(a'_n)_{n \in \mathbb{N}}$ e che $(b_n)_{n \in \mathbb{N}}$ sia equiconvergente a $(b'_n)_{n \in \mathbb{N}}$. Dobbiamo provare che $(a_n \cdot b_n)_{n \in \mathbb{N}}$ è equiconvergente a $(a'_n \cdot b'_n)_{n \in \mathbb{N}}$. Posto $\delta_n = a'_n - a_n$ e $\gamma_n = b'_n - b_n$, le due successioni $(\delta_n)_{n \in \mathbb{N}}$ e $(\gamma_n)_{n \in \mathbb{N}}$ convergono a zero. Inoltre, poichè

$$(a'_n \cdot b'_n)_{n \in \mathbb{N}} = ((a_n + \delta_n) \cdot (b_n + \gamma_n))_{n \in \mathbb{N}}$$

¹⁹ Come osservato nell'appendice, un modo per provare che un anello non è un campo è provare che ammette divisori dello zero. Nel nostro caso la coppia costituita da $(a_n)_{n \in \mathbb{N}}$ e dalla successione $(b_n)_{n \in \mathbb{N}}$ definita dal porre $b_n = 0$ se n è pari e $b_n = 1/n$ se n è dispari, è una coppia di divisori dello zero.

²⁰ Invece di riferirci alla nozione di equi-convergenza possiamo anche riferirci alla teoria degli ideali in un anello. Infatti indichiamo con I l'insieme delle successioni di Cauchy $(a_n)_{n \in \mathbb{N}}$ tali che $\lim_{n \rightarrow \infty} |a_n| = 0$. Allora I costituisce un ideale dell'anello $(Ch, +, \cdot, \underline{0}, \underline{1})$. Si prova che tale ideale è massimale. Inoltre due successioni $(a_n)_{n \in \mathbb{N}}$ e $(b_n)_{n \in \mathbb{N}}$ sono equiconvergenti se e solo se la loro differenza appartiene ad I . Pertanto possiamo definire il campo dei reali come il quoziente di $(Ch, +, \cdot, \underline{0}, \underline{1})$ modulo l'ideale massimale I .

$$\begin{aligned}
&= (a_n \cdot b_n + a_n \gamma_n + \delta_n \cdot b_n + \delta_n \gamma_n)_{n \in N} \\
&= (a_n \cdot b_n)_{n \in N} + (a_n \gamma_n)_{n \in N} + (\delta_n \cdot b_n)_{n \in N} + (\delta_n \gamma_n)_{n \in N}
\end{aligned}$$

e poiché $(a_n \gamma_n)_{n \in N}$, $(\delta_n \cdot b_n)_{n \in N}$ e $(\delta_n \gamma_n)_{n \in N}$ convergono a zero, $(a_n \cdot b_n)_{n \in N}$ risulta equiconvergente a $(a'_n \cdot b'_n)_{n \in N}$. Allo stesso modo si dimostra la compatibilità della equiconvergenza rispetto alla addizione.

Proviamo ora che $(Ch/\equiv, +, \cdot, [0], [1])$ è un campo. A tale scopo sia $[(a_n)_{n \in N}]$ diversa dalla classe nulla. Allora $(a_n)_{n \in N}$ non è equiconvergente alla successione costantemente uguale a zero e quindi non converge a zero. Ne segue che, per il teorema della permanenza del segno, esiste $m \in N$ tale che $a_n \neq 0$ per ogni $n \geq m$. Sia $b_n = 1/a_n$ per ogni $n \geq m$ e $b_n = 1$ altrimenti. Allora risulta che la successione $(a_n \cdot b_n)_{n \in N}$ è costantemente uguale ad 1 tranne per un numero finito di elementi ed è quindi equiconvergente a $\underline{1}$. In definitiva $[(a_n)_{n \in N}] \cdot [(b_n)_{n \in N}] = [\underline{1}]$. $\square$

Definizione 10.6. In $Ch/\equiv$ definiamo la relazione $<$ ponendo $[(a_n)_{n \in N}] < [(b_n)_{n \in N}]$ se esiste m tale che $a_n < b_n$ per ogni $n \geq m$. Definiamo la relazione $\leq$ ponendo $[(a_n)_{n \in N}] \leq [(b_n)_{n \in N}]$ se $[(a_n)_{n \in N}] < [(b_n)_{n \in N}]$ oppure $[(a_n)_{n \in N}] = [(b_n)_{n \in N}]$.

Omettiamo la dimostrazione del seguente fondamentale teorema. Per la nozione di campo ordinato completo si veda nel capitolo 5 il paragrafo sull'approccio assiomatico ai numeri reali.

Teorema 10.7. La struttura $(Ch/\equiv, +, \cdot, \leq, [0], [1])$ è un campo ordinato completo che chiamiamo *campo dei numeri reali*.

Concludiamo questo paragrafo osservando che si pone il problema di come si possa rappresentare un numero reale. Ora è ovviamente poco pratico considerare una classe completa di equivalenza nell'insieme delle successioni di Cauchy. Allora si deve trovare un modo per scegliere all'interno di ciascuna classe un elemento che la rappresenti. La rappresentazione in base 10 dei numeri reali consiste nel considerare all'interno di ogni classe una particolare successione di Cauchy è precisamente una serie di potenze di base 10. Infatti quando indichiamo con $a_n a_{n-1} \dots a_0, c_1 c_2 \dots$ intendiamo infatti il numero reale corrispondente alla serie

$$a_n 10^n + \dots + a_0 10^0 + c_1 10^{-1} + \dots$$

che risulta essere una successione di Cauchy. Naturalmente sarebbe necessario dimostrare che ogni numero reale si può rappresentare in questo modo, cioè che ogni successione di Cauchy è equiconvergente ad una serie di potenze di base 10.

Problema: Dire quale è il significato dell'asserzione “il numero 0,399... è uguale al numero 0.4000...” e quale sarebbe la cosa da dimostrare per controllare che tale affermazione è vera.

11. Un percorso diverso: essere quasi uguali

Il campo dei numeri reali non è l'unica base per la costruzione della matematica. Infatti esistono altri modi di estendere il campo dei razionali ottenendo strutture che risultano più ricche ed interessanti del campo dei reali per il fatto di contenere “infinitesimi” ed “infiniti”.²¹ Partiamo ancora una volta dall'anello $(Q^N, +, \cdot, 0, 1)$ delle successioni di razionali. Tuttavia questa volta:

1. non ci limitiamo alle successioni di Cauchy ma consideriamo la classe di tutte le successioni
2. introduciamo una relazione di congruenza diversa dalla equi-convergenza.

L'idea su cui ci baseremo è che si possono identificare due successioni che sono uguali in un insieme di indici che ci sembra “sufficientemente grande” o, se si vuole, due successioni che sono “uguali quasi ovunque”. Ricordiamo che in matematica esistono due nozioni di “quasi ovunque” che sembrano adeguate. Infatti viene detto che una proprietà P vale quasi ovunque quando:

- “è sufficiente che P valga ovunque tranne che per un insieme finito di elementi”

oppure

- “è sufficiente che P valga ovunque tranne che per un insieme di elementi di misura nulla”.

Se indichiamo con U la classe degli insiemi co-finiti (cioè complementi di finiti), allora nel primo caso possiamo dire che P vale quasi ovunque se l'insieme degli elementi in cui P vale appartiene ad U . Similmente se U denota la classe degli insiemi che sono complementi di insiemi di misura nulla, allora anche nel

²¹ Per la comprensione di questo paragrafo si suggerisce di leggere prima il paragrafo 5 del capitolo 5 in cui si parla dell'approccio assiomatico ai numeri reali.

secondo caso possiamo dire che l'insieme degli elementi in cui vale P appartiene ad U . Inoltre in entrambi i casi la classe U soddisfa le seguenti proprietà.

- i) $X \in U$ e $Y \in U \Rightarrow X \cap Y \in U$
- ii) $X \in U$ e $Y \supseteq X \Rightarrow Y \in U$.

Tali proprietà suggeriscono di estendere le due nozioni di “quasi ovunque” nel modo più generale possibile.

Definizione 11.1. Dato un insieme S chiamiamo *filtro* una classe U non vuota di sottoinsiemi di S tale che

- i) $X \in U$ e $Y \in U \Rightarrow X \cap Y \in U$
- ii) $X \in U$ e $Y \supseteq X \Rightarrow Y \in U$

Diciamo che una proprietà P vale *quasi ovunque* in S rispetto ad U se $\{x \in S : x \text{ soddisfa } P\} \in U$.

$P(S)$ è un filtro che viene chiamato *improprio* (il più grande). Dalla proprietà ii) segue che un filtro è improprio se e solo se contiene l'insieme vuoto. Pertanto, per i), in un filtro proprio non può capitare che esistano due sottoinsiemi disgiunti. Un altro esempio di filtro è dato da $\{S\}$ (il più piccolo).

Esercizio. Fissato un sottoinsieme A di S provare che la classe $U = \{X \in P(S) : X \supseteq A\}$ degli insiemi che contengono A è un filtro (tale tipo di filtro viene detto *principale*). In tale caso una proprietà vale quasi ovunque se vale almeno per tutti gli elementi di A . In un certo senso questo vuol dire che si considerano importanti solo gli elementi di A e trascurabili gli altri. Detto in altre parole, “vero quasi ovunque” significa “vero per tutti gli elementi importanti”.

Esercizio. Provare che la classe degli intorni di un punto x nel piano (più in generale in uno spazio topologico) costituisce un filtro. In tale caso una proprietà P vale quasi ovunque se P vale in un intorno di x .

Definizione 11.2. Detto U un filtro nell'insieme N , definiamo in Q^N la relazione binaria $\equiv$ ponendo

$$(a_n)_{n \in N} \equiv (b_n)_{n \in N} \Leftrightarrow a_n = b_n \text{ quasi ovunque rispetto ad } U.$$

In altre parole poniamo $(a_n)_{n \in N} \equiv (b_n)_{n \in N}$ se e solo se $\{n \in N : a_n = b_n\} \in U$. Se U è il filtro improprio allora due successioni sono sempre equivalenti. Se U è il filtro $\{N\}$ allora due successioni sono equivalenti solo se coincidono. Se U è il filtro generato dall'insieme $X \subseteq N$, allora due successioni sono equivalenti solo se hanno gli stessi valori in X .

Teorema 11.3. La relazione $\equiv$ è una congruenza nella struttura $(Q^N, +, \cdot, \underline{0}, \underline{1})$ ed il relativo quoziente $(Q^N/\equiv, +, \cdot, [\underline{0}], [\underline{1}])$ è un anello.

Dim. Proviamo che $\equiv$ è una equivalenza. La proprietà riflessiva e simmetrica sono “ereditate” dalle corrispondenti proprietà dell'uguaglianza. Infatti, poiché $\{n \in N : a_n = a_n\} = N \in U$, risulta che $(a_n)_{n \in N} \equiv (a_n)_{n \in N}$. Assumiamo che $(a_n)_{n \in N} \equiv (b_n)_{n \in N}$ e quindi che $\{n \in N : a_n = b_n\} \in U$. Allora, poiché $\{n \in N : b_n = a_n\} = \{n \in N : a_n = b_n\}$, risulta che $(b_n)_{n \in N} \equiv (a_n)_{n \in N}$. Ciò prova la proprietà simmetrica. Per provare la proprietà transitiva, supponiamo che $(a_n)_{n \in N} \equiv (b_n)_{n \in N}$ e $(b_n)_{n \in N} \equiv (c_n)_{n \in N}$, cioè che $\{n \in N : a_n = b_n\} \in U$ e $\{n \in N : b_n = c_n\} \in U$. In tale caso, poiché

$$\{n \in N : a_n = c_n\} \supseteq \{n \in N : a_n = b_n\} \cap \{n \in N : b_n = c_n\}$$

risulta che $\{n \in N : a_n = c_n\} \in U$ e quindi che $(a_n)_{n \in N} \equiv (c_n)_{n \in N}$.

Per provare che $\equiv$ è una congruenza, supponiamo che $(a_n)_{n \in N} \equiv (a'_n)_{n \in N}$ e $(b_n)_{n \in N} \equiv (b'_n)_{n \in N}$, cioè che

$$\{n \in N : a_n = a'_n\} \in U \text{ e } \{n \in N : b_n = b'_n\} \in U.$$

In tale caso, poiché

$$\{n \in N : a_n + b_n = a'_n + b'_n\} \supseteq \{n \in N : a_n = a'_n\} \cap \{n \in N : b_n = b'_n\}$$

possiamo concludere che $(a_n)_{n \in N} + (b_n)_{n \in N} \equiv (a'_n)_{n \in N} + (b'_n)_{n \in N}$. Esattamente nello stesso modo è possibile provare che $\equiv$ è compatibile con il prodotto.

Infine per provare che $(Q^N/\equiv, +, \cdot, [\underline{0}], [\underline{1}])$ è un anello osserviamo che il passaggio a quoziente di una struttura conserva tutte le proprietà che si possono esprimere tramite equazioni. Poiché gli assiomi di anello sono espressi tutti tramite equazioni, $(Q^N/\equiv, +, \cdot, [\underline{0}], [\underline{1}])$ essendo quoziente di un anello è un anello. $\square$

Se U è il filtro $\{N\}$ allora $Q^N/\equiv$ coincide con l'anello Q^N . Se U è il filtro principale generato da un singleton allora è evidente che $Q^N/\equiv$ è isomorfo a Q .

12. I razionali non-standard

L'anello $(Q^N/\equiv, +, \cdot, [0], [1])$ non è un campo, in generale. Ad esempio consideriamo il filtro dei cofiniti e le due successioni $(a_n)_{n \in N}$ e $(b_n)_{n \in N}$ definite dall'essere

$$\begin{aligned} a_n &= 1/n \text{ se } n \text{ è pari e } a_n = 0 \text{ se } n \text{ è dispari} \\ b_n &= 0 \text{ se } n \text{ è pari e } b_n = 1/n \text{ se } n \text{ è dispari,} \end{aligned}$$

Allora

$$[(a_n)_{n \in N}] \cdot [(b_n)_{n \in N}] = [(a_n)_{n \in N} \cdot (b_n)_{n \in N}] = [(a_n \cdot b_n)_{n \in N}] = [0]$$

con $[(a_n)_{n \in N}] \neq [0]$ e $[(b_n)_{n \in N}] \neq [0]$. Per evitare un tale problema dobbiamo considerare un filtro che contenga o l'insieme dei numeri pari o l'insieme dei numeri dispari in modo che risulti o che $[(a_n)_{n \in N}] = [0]$ oppure che $[(b_n)_{n \in N}] \neq [0]$. Più in generale, poiché possiamo definire due successioni analoghe per ogni sottoinsieme X di N ponendo

$$\begin{aligned} a_n &= 1/n \text{ se } n \in X \text{ e } a_n = 0 \text{ se } n \notin X \\ b_n &= 0 \text{ se } n \notin X \text{ e } b_n = 1/n \text{ se } n \in X, \end{aligned}$$

è evidente che per evitare che ci siano divisori dello zero dobbiamo supporre che U sia un filtro tale che per ogni sottoinsieme X risulti che $X \in U$ oppure $-X \in U$. Ciò conduce alla seguente definizione.

Definizione 12.1. Diciamo che un filtro proprio U è un *ultrafiltro* se per ogni $X \in P(S)$,

$$iii) \quad X \in U \text{ oppure } -X \in U.$$

Proposizione 12.2. La condizione *iii*) equivale a

$$iii^*) \quad X \cup Y \in U \Rightarrow X \in U \text{ oppure } Y \in U.$$

Dim. Supponiamo che valga *iii*) e che $X \cup Y \in U$, allora nel caso in cui $X \notin U$, poiché $-X \in U$ deve risultare che $Y - X = (X \cup Y) \cap -X \in U$. Pertanto Y contenendo $Y - X$, appartiene ad U . Ciò prova *iii**). Supponiamo che valga *iii**), allora, essendo $X \cup -X = S \in U$ risulta che o $X \in U$ oppure $-X \in U$. Ciò prova *iii*). $\square$

Proposizione 12.3. Se U è un ultrafiltro le seguenti asserzioni sono equivalenti:

- a) U è principale
- b) U è generato da un singleton
- c) U contiene un insieme finito.

Ne segue che un ultrafiltro non principale contiene tutti i cofiniti.

Dim. $a) \Rightarrow b)$ Sia U un ultrafiltro principale generato da un insieme X e supponiamo che X non sia un singleton. Allora X si scomporrebbe in due sottoinsiemi propri e disgiunti X_1 e X_2 . Per $iii)$ uno di tali due insiemi deve appartenere ad U in contrasto con l'ipotesi che tutti gli elementi di U contengono X .

$b) \Rightarrow c)$ Evidente.

$c) \Rightarrow a)$ Supponiamo che U sia un ultrafiltro contenente un insieme finito e sia X un insieme che abbia tra tutti i finiti in U il numero minimo di elementi. Allora X non può spezzarsi in due sottoinsiemi propri poiché uno di questi sarebbe un finito più piccolo appartenente ad U . Quindi X è un singleton del tipo $\{x\}$. Poiché l'intersezione di due elementi di un filtro proprio è sempre non vuota, un insieme X appartiene ad U se e solo se contiene x .

Infine se l'ultrafiltro U non è principale allora non può contenere nessun insieme finito. Questo comporta che deve contenere i complementi di tutti gli insiemi finiti, cioè deve contenere tutti i cofiniti. $\square$

A noi interessano ultrafiltri che non sono principali poiché per quelli principali l'ultrapotenza di Q coincide con Q . Non è facile dare esempi intuitivi di ultrafiltri che non siano principali. Infatti la dimostrazione della loro esistenza viene effettuata mediante l'assioma della scelta e non tramite l'esplicita esibizione di un esempio. Enunciamo solo il teorema che esprime una tale esistenza.

Teorema 12.4. Esiste un ultrafiltro che non è principale.

Teorema 12.5. Sia U un ultrafiltro, allora il quoziente di $(Q^N, +, \cdot, \underline{0}, \underline{1})$ modulo $\equiv$ è un campo che estende il campo dei razionali e che denotiamo con Q^* . Se U non è principale, allora diciamo che Q^* è un *campo di numeri razionali non standard*.

Dim. Per provare che Q^* è un campo supponiamo che $[(a_n)_{n \in N}]$ sia un elemento di Q^* diverso da zero. Allora, poiché lo zero di Q^* è la classe determinata dalla successione costantemente uguale a zero, sappiamo che $\{n \in N : a_n = 0\} \notin U$. Poiché per ipotesi U è un ultrafiltro, ciò comporta che $\{n \in N : a_n \neq 0\} \in U$. Definiamo

la successione $(b_n)_{n \in \mathbb{N}}$ ponendo $b_n = 1/a_n$ se $a_n \neq 0$ e $b_n = 1$ altrimenti. Allora, poiché

$$\{n \in \mathbb{N} : a_n \cdot b_n = 1\} \supseteq \{n \in \mathbb{N} : a_n \neq 0\} \in U,$$

risulta che $\{n \in \mathbb{N} : a_n \cdot b_n = 1\} \in U$ e quindi che

$$[(a_n)_{n \in \mathbb{N}}] \cdot [(b_n)_{n \in \mathbb{N}}] = [(a_n b_n)_{n \in \mathbb{N}}] = [\underline{1}].$$

Ciò prova che $[(a_n)_{n \in \mathbb{N}}]$ è invertibile.

Per mostrare che Q^* estende Q è sufficiente considerare l'applicazione che associa ad ogni razionale r la classe di equivalenza $[(r_n)_{n \in \mathbb{N}}]$ dove $(r_n)_{n \in \mathbb{N}}$ è la successione costantemente uguale ad r . Tale applicazione risulta essere una immersione. $\square$

Sia il campo dei razionali che quello dei reali risultano essere un *campo ordinato*, (si veda la definizione in Appendice). Possiamo anche in Q^* definire una relazione d'ordine trasmettendo la relazione di ordinamento tra successioni definita in $Q^{\mathbb{N}}$ al quoziente Q^* . Purtroppo mentre sappiamo bene che cosa è una congruenza in una struttura algebrica e come effettuare il quoziente di tale struttura, nel caso siano coinvolte relazioni il discorso è alquanto più problematico. Un modo semplice è tuttavia quello di ricondurre la nozione di ordine ad una operazione come avviene in teoria dei reticoli. Infatti in un reticolo viene definita una relazione d'ordine ponendo $x \leq y$ se e solo se $x \wedge y = x$. Allora consideriamo Q come reticolo rispetto alle operazioni *max* e *min* ed estendiamo tali operazioni al prodotto diretto $Q^{\mathbb{N}}$. Si ottiene il reticolo $(Q^{\mathbb{N}}, \min, \max)$. Per tale reticolo risulta:

Proposizione 12.6. Dato un ultrafiltro, la relativa relazione di equivalenza $\equiv$ è una congruenza nel reticolo $(Q^{\mathbb{N}}, \min, \max)$ ed il quoziente $(Q^{\mathbb{N}}/\equiv, \min, \max)$ di tale reticolo è ancora un reticolo. Pertanto è definita in $Q^{\mathbb{N}}/\equiv$ una relazione d'ordine $\leq$ ponendo $[(a_n)_{n \in \mathbb{N}}] \leq [(b_n)_{n \in \mathbb{N}}]$ se e solo se $[(a_n)_{n \in \mathbb{N}}] \wedge [(b_n)_{n \in \mathbb{N}}] = [(a_n)_{n \in \mathbb{N}}]$. Tale relazione d'ordine può essere definita anche ponendo:

$$[(a_n)_{n \in \mathbb{N}}] \leq [(b_n)_{n \in \mathbb{N}}] \Leftrightarrow a_n \leq b_n \text{ quasi ovunque.}$$

Dim. La prima parte della proposizione si prova come nella dimostrazione del Teorema 11.3. D'altra parte

$$\begin{aligned} [(a_n)_{n \in \mathbb{N}}] \leq [(b_n)_{n \in \mathbb{N}}] &\Leftrightarrow \min\{a_n, b_n\} = a_n \text{ quasi ovunque} \\ &\Leftrightarrow a_n \leq b_n \text{ quasi ovunque.} \end{aligned} \quad \square$$

Una aspetto particolarmente interessante di un campo dei razionali non standard è che non è archimedeo.

Teorema 12.7. Se l'ultrafiltro U contiene il filtro dei cofiniti allora il corrispondente campo Q^* dei numeri razionali non standard è un campo ordinato non archimedeo rispetto all'ordinamento che abbiamo ora definito.

Dim. Consideriamo la successione n^2 ed il corrispondente numero iperreale $[(n^2)_{n \in N}]$. Allora preso un qualunque intero p risulta che $\{n \in N : n^2 < p\}$ è finito e quindi che $\{n \in N : n^2 \geq p\}$ è cofinito. Poiché abbiamo supposto che U contiene tutti i cofiniti, tale insieme appartiene ad U . Questo prova che $[(n^2)_{n \in N}] \geq p \cdot [1]$ qualunque sia p , cioè che non esiste nessun multiplo di $[1]$ capace di superare $[(n^2)_{n \in N}]$. $\square$

Il fatto che Q^* non sia archimedeo comporta che non può essere isomorfo al campo dei numeri reali (che invece sappiamo essere un campo completo e quindi archimedeo). D'altra parte nella seguente proposizione mostriamo che in Q^* non è possibile nemmeno l'estrazione della radice quadrata.

Proposizione 12.8. In Q^* l'equazione $x^2 = 2$ non ammette soluzioni, cioè non esiste la radice di 2.

Dim. Supponiamo che esista $[(a_n)_{n \in N}]$ in Q^* tale che $[(a_n)_{n \in N}] \cdot [(a_n)_{n \in N}] = [2]$. Allora dovrebbe essere che $a_n^2 = 2$ quasi ovunque cosa impossibile poiché nessun numero razionale verifica tale uguaglianza (la radice quadrata di 2 non è razionale). $\square$

Sembrerebbe quindi che Q^* non sia l'ambiente giusto per potere costruire tutta la matematica come invece sembra essere il campo dei reali. Tuttavia vale il seguente teorema

Teorema 12.9. Esiste un x in Q^* tale che x^2 è *infinitamente vicino* a 2, cioè tale che la differenza $x^2 - 2$ è un infinitesimo. Più in generale, per ogni numero reale x esiste un elemento x' di Q^* infinitamente vicino ad x .

Non andiamo oltre a tali questioni perché lo scopo di questo paragrafo è solo quello di fare intravedere un universo nuovo ed

interessante che è quello dell'*analisi non standard*. Ci limitiamo ad osservare che con la stessa tecnica con cui abbiamo costruito il campo Q^* dei razionali non standard possiamo costruire il campo dei reali non standard. E' sufficiente partire dalle successioni di numeri reali invece che dalle successioni di razionali.

Teorema 12.10. Sia U un ultrafiltro che contiene il filtro dei cofiniti e definiamo in R^N la relazione $\equiv$ ponendo

$$(a_n)_{n \in N} \equiv (b_n)_{n \in N} \Leftrightarrow \{n \in N : a_n = b_n\} \in U.$$

Allora $\equiv$ è una congruenza nella struttura $(R^N, +, \cdot, \underline{0}, \underline{1})$. Il relativo quoziente $(Q^N/\equiv, +, \cdot, [\underline{0}], [\underline{1}])$ è un campo non archimedeo rispetto all'ordinamento definito come in proposizione 12.4. Tale campo prende il nome di *campo dei numeri reali non-standard*.

L'approccio all'analisi matematica che utilizza un campo di numeri reali non standard prende il nome di *Analisi non standard*. In un certo senso l'analisi non standard è un modo puramente algebrico di trattare il calcolo differenziale. Infatti la presenza di infiniti ed infinitesimi permette, per fare un esempio, di definire il limite di una funzione $f(x)$ per x che tende all'infinito direttamente come il valore di f calcolato in un numero infinito. Un integrale definito è effettivamente una sorta di somma infinitaria, la derivata è il rapporto tra due infinitesimi e così via. Ciò spesso fornisce metodi eleganti e naturali per dimostrare teoremi in analisi matematica.

LETTURA

Cesare Zavattini, *Gara di Matematica*: da *I tre libri, Parliamo tanto di me* - Bompiani - cap. XVI pag. 48,49,50.²²

E' un ricordo della mia infanzia. Abitavo a Gottinga nel dicembre del milleottocentosettanta. Mio padre ed io giungemmo all'Accademia quando il presidente Maust stava cominciando l'appello dei partecipanti alla *Gara Mondiale di Matematica*. Subito babbo andò a mettersi fra gli iscritti dopo avermi affidato alla signora Katten, amica di famiglia. Seppi da lei che il colpo del cannone di Pombo, il bidello, avrebbe segnato l'inizio della storica contesa. La signora Katten mi raccontò un episodio, ignoto ai più, intorno all'attività di Pombo. Costui sparava da trent'anni un colpo di cannone per annunciare il mezzogiorno preciso. Una volta se n'era dimenticato. Il dì appresso, allora, aveva sparato il colpo del giorno prima, e così di seguito fino a quel venerdì del milleottocentosettanta. Nessuno a Gottinga si era mai accorto che Pombo sparava il colpo del giorno avanti.²³ Esauriti i preliminari, la gara ebbe inizio alla presenza del principe Ottone e di un ragguardevole gruppo di intellettuali.

“Uno, due, tre, quattro, cinque...”

Nella sala si udivano soltanto le voci dei gareggianti. Alle diciassette circa, avevano superato il ventesimo migliaio. Il pubblico si appassionava alla nobile contesa e i commenti si intrecciavano. Alle diciannove, Alain, della Sorbona, si accasciò sfinito.

Alle venti, i superstiti erano sette.

“36767, 36768, 36769, 36770...”

Alle ventuno Pombo accese i lampioni. Gli spettatori ne approfittarono per mangiare le provviste portate da casa.

“40719, 40720, 40721...”

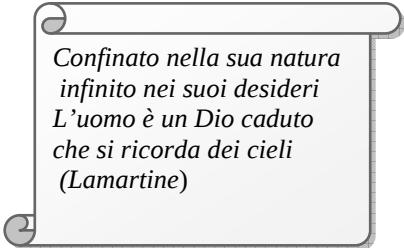
Io guardavo mio padre, madido di sudore, ma tenace. La signora Katten accarezzandomi i capelli ripeteva come un ritornello: 'Che bravo babbo hai,' e a me non pareva neppure di avere fame. Alle ventidue precise avvenne il primo colpo di scena: l'algebrista Pull scattò:

²² Esponiamo un breve racconto di Zavattini che in un certo senso parla dell'infinito e dell'operatore di successore (e quindi della natura di una terna di Peano).

²³ Tale episodio può essere visto come il fatto che per un cardinale infinito x accade che $x+1 = x$ (si veda il capitolo successivo sulla cardinalità).

CAPITOLO 4

GLI INSIEMI: CREDERE NELL'INFINITO¹



*Confinato nella sua natura
infinito nei suoi desideri
L'uomo è un Dio caduto
che si ricorda dei cieli
(Lamartine)*

1. Il prezzo dell'arimetizzazione: l'infinito attuale

Nel capitolo precedente abbiamo mostrato come sia possibile una arimetizzazione di tutta la matematica che rende possibile svincolare il discorso matematico dall'intuizione del continuo geometrico. Tuttavia per portare avanti questa arimetizzazione è necessario “pagare un prezzo”. Questo prezzo è l'accettazione dell'infinito potenziale per potere definire gli interi ed i razionali e, come vedremo, l'accettazione dell'infinito attuale per potere definire i reali. Naturalmente l'insieme degli elementi di una terna di Peano è attualmente infinito (anzi l'accettazione dell'esistenza di una terna di Peano equivale all'accettazione di un insieme infinito). Tuttavia i singoli numeri interi sono oggetti finiti e quando effettuiamo i nostri calcoli ci serviamo ogni volta di una quantità finita di numeri. La stessa cosa può essere detta per i numeri relativi ed i razionali anche se apparentemente un razionale o un relativo è una classe attualmente infinita di coppie (e quindi un oggetto infinito). Ad esempio il numero razionale che indichiamo con $2/3$ è la classe $[(2,3)]$ che è l'insieme infinito di tutte le coppie equivalenti alla coppia $(2,3)$, cioè l'insieme $\{(2n,3n): n \in \mathbb{N}\}$. Tuttavia possiamo definire gli interi relativi ed i razionali anche senza coinvolgere la nozione di classe completa

¹ In questo capitolo parleremo della teoria “ingenua” degli insiemi. Questo significa che gli insiemi verranno introdotti in modo completamente intuitivo. In realtà, come vedremo nel prossimo capitolo, esistono molti paradossi della teoria degli insiemi che mostrano come un approccio informale presenti molti problemi e come sia necessario abbandonare la fede (ingenua) verso l'accettazione di insiemi infiniti.

di equivalenza. Basta fissare in ogni classe un elemento rappresentativo e lavorare solo sugli elementi rappresentativi.

Definizione 1.1. Chiamiamo in *forma normale*² ogni coppia di numeri naturali del tipo $(p,0)$, che possiamo indicare con $+p$, oppure del tipo $(0,p)$, che possiamo indicare con $-p$. Indichiamo con Z_n l'insieme delle coppie in forma normale.

Riferendoci alla relazione di equivalenza che abbiamo definito per introdurre l'anello dei numeri relativi, si chiama *riduzione a forma normale* il calcolo che partendo da una coppia (m,n) permette di ottenere la coppia equivalente in forma normale. Precisamente (m,n) si riduce ad $(m-n,0)$ se $m \geq n$ ed a $(0,n-m)$ se $m < n$. La somma ed il prodotto di due forme normali non è in generale una forma normale. Tuttavia potremmo definire in Z_n una struttura algebrica al modo seguente.

Definizione 1.2. Chiamiamo *anello degli interi relativi* la struttura $(Z_n, \oplus, \otimes, (0,0), (1,0))$ in cui le operazioni sono definite ponendo:

- $x \oplus y$ uguale alla riduzione a forma normale di $x+y$
- $x \otimes y$ uguale alla riduzione a forma normale di $x \cdot y$.

Lo stesso discorso può essere fatto per i numeri razionali. Infatti ogni classe $[(p,q)]$ in Q può essere rappresentata in un solo modo da una coppia (p,q) con p e q primi tra loro.

Definizione 1.3. Chiamiamo in *forma normale* una coppia (p,q) di interi relativi con $q > 0$ e p e q primi tra loro. Indichiamo con Q_n l'insieme delle coppie in forma normale.

Ogni coppia (p,q) può essere ridotta in forma normale dividendo per tutti gli eventuali fattori comuni.

Definizione 1.4. Chiamiamo *campo dei numeri razionali* la struttura $(Q_n, \oplus, \otimes, (0,1), (1,1))$ in cui,

- $x \oplus y$ è uguale alla riduzione a forma normale di $x+y$

² La nozione di *forma normale*, e lo studio connesso dei metodi di *riduzione a forma normale*, gioca un ruolo fondamentale in molti campi della matematica.

-
- $x \otimes y$ è uguale alla riduzione a forma normale di $x \cdot y$.

In definitiva quando si maneggiano gli interi relativi o i razionali di fatto è possibile maneggiare solo le forme normali che sono oggetti finiti.

Questo modo di procedere è molto generale ed in un certo senso può sostituire l'usuale definizione di quoziente di una struttura algebrica modulo una data congruenza. Infatti, data una congruenza (ma è sufficiente anche una qualunque equivalenza) in una data struttura algebrica, invece di lavorare sulle classi di equivalenza definendo su tali classi le relative operazioni, è possibile:

- individuare all'interno di ogni classe un particolare elemento che viene detto *in forma normale*
- individuare un *procedimento di riduzione a forma normale* che permetta, dato un elemento x di trovare un elemento x' equivalente ad x e ridotto a forma normale
- considerare solo gli elementi in forma normale
- effettuare le operazioni sulle forme normali e poi ridurre il risultato ottenuto in forma normale.

D'altra parte questo è il modo effettivo come i matematici trattano i numeri.³ Lo stesso modo è adottato dai sistemi di calcolo simbolico come *Mathematica* che, d'altra parte, sono sistemi di intelligenza artificiale che simulano il comportamento di un matematico. Ad esempio, dovendo effettuare la somma tra $6/9$ e $10/12$, *Mathematica* riduce tali coppie a forma normale ottenendo, come rappresentanti delle corrispondenti classi, le coppie $2/3$ e $5/6$. Poi effettua l'operazione di somma di due coppie ottenendo prima $(2 \cdot 6 + 3 \cdot 5) / 3 \cdot 6$, e quindi, dopo avere effettuato le operazioni tra interi, $27/18$. Infine riduce a forma normale tale coppia ed ottiene come risultato $3/2$. Allo stesso modo *Mathematica* riesce a trattare una larga parte della matematica.

³ Anche negli interi modulo un intero m in effetti si sceglie come rappresentativo di una classe il numero positivo più piccolo (che risulta minore di $m-1$) ci si riferisce alle "forme normali" $0, 1, \dots, m-1$ e non alle rispettive classi. Ad esempio, nel caso $m = 5$, la somma di 4 e 3 viene fatta prima addizionando i numeri, si ottiene 7 e poi, riducendo a forma normale, si ottiene 2 . In breve $4+3 = 2$.

Invece il passaggio dai razionali ai reali crea un coinvolgimento inevitabile dell'infinito attuale e questo poiché, qualunque sia il modo con cui si sono costruiti i reali,

ogni numero reale è un oggetto infinito.

Per convincersi di questo, riferendoci al metodo delle sezioni osserviamo che una sezione è costituita da due insiemi infiniti di razionali. Se invece ci riferiamo al metodo delle successioni di Cauchy, possiamo osservare che una successione di Cauchy è un oggetto infinito e che una classe completa di equivalenza di successioni di Cauchy è un oggetto infinito.⁴

Ora abbiamo già visto come, da Aristotele in poi, fosse netto nel mondo greco il rifiuto dell'infinito attuale. Tale rifiuto fu successivamente condiviso da quasi tutta la cultura occidentale fino alla fine dell'ottocento. Allo stesso tempo l'impetuoso sviluppo dell'analisi matematica dal 1600 in poi aveva fatto sì che l'uso dei metodi infinitari fosse sempre più una cosa inevitabile. L'alternativa che spesso si presentava agli scienziati dell'epoca era tra lo sterile rigore della geometria euclidea e l'uso spregiudicato dei nuovi metodi infinitari del calcolo differenziale ed integrale. A sua volta l'accettazione dei metodi infinitari rendeva l'ambito dell'algebra e della geometria greca troppo ristretto per la matematica moderna. Infatti, per fare un esempio, è chiaro che l'algebra e la geometria suggerivano e permettevano solo lo studio delle funzioni elementari, cioè quelle definibili geometricamente (come le funzioni trigonometriche), quelle definibili algebricamente (come i polinomi o le funzioni razionali) e quelle ottenibili per composizione da queste. Invece lo sviluppo delle serie trigonometriche determinò un enorme allargamento del campo delle funzioni note. Ci si accorse che, a partire dalle note funzioni trigonometriche, ed operando con somme infinite, era possibile pervenire a nuove funzioni che non erano definibili né per via geometrica né per via algebrica. Era quindi necessario dare una definizione più astratta e generale del concetto di funzione.

La teoria degli insiemi proposta da Cantor alla fine del ottocento fornirà lo strumento adatto allo scopo.

⁴ Precisamente una successione ha per definizione la potenza del numerabile, un numero reale, visto come classe completa di successioni equiconvergenti, ha la potenza del continuo (provare a trovare una dimostrazione).

2. Ma questi insiemi sono poi veramente una novità ?

Le generazioni seguenti considereranno la teoria degli insiemi come una malattia da cui si è guariti (Henri Poincaré, 1908).

Cantor non è certo stato il primo ad utilizzare concetti come quelli di insieme, classe, collezione, concetti questi che si sono sempre adoperati sia nel discorso scientifico che nel linguaggio comune. In effetti ogni volta che si considera una proprietà apparire naturale considerare la collezione di tutti gli oggetti che verificano tale proprietà (la sua estensione). Ad esempio:

- alla proprietà di essere mammifero corrisponde "la classe di tutti i mammiferi",
- alla proprietà di riprodursi tramite uova corrisponde "la classe degli ovipari",
- etc.

Ma allora se la nozione di classe è sempre esistita:

- *perché si dice che Cantor fu l'inventore della teoria degli insiemi ?*
- *in che cosa consiste la novità della sua proposta ?*
- *può esistere un mondo senza insiemi ?*

Per cercare di capire come vanno le cose, consideriamo alcune frasi del linguaggio comune, ad esempio le frasi:

- a) *la rosa che è nel vaso è rossa;*
- b) *le rose del vaso sono rosse;*
- c) *le rose del vaso sono dodici.*

La prima frase è del tipo usuale, vi è un soggetto (la rosa che è nel vaso) ed un predicato (essere rosso).⁵ Le altre due frasi coinvolgono invece collezioni di elementi (di sostanze) come testimonia l'uso del plurale "le rose", ma, a bene osservare, ciò avviene in due modi totalmente diversi. Infatti nella frase b) il predicato "essere rosso" non si riferisce all'insieme delle rose del fascio: non si è mai visto un insieme rosso. Il predicato si riferisce in realtà ai singoli elementi di tale insieme, cioè a ciascuna

⁵ Ad una struttura della frase di tale tipo corrisponde ad una concezione del mondo secondo cui da un lato vi sono delle "sostanze" e dall'altro delle "proprietà" di cui tali sostanze possono godere o meno. Questo modo di vedere è espresso in netta da Aristotele e gli impedirà di capire il significato e l'importanza delle relazioni binarie.

rosa nel vaso. La *b*) è in realtà un modo abbreviato per affermare che:

b') *ciascuna rosa che è nel vaso è rossa.*

Volendo utilizzare a tutti i costi gli insiemi, possiamo anche riscrivere la *b*) dicendo che:

b'') *l'insieme delle rose del vaso è contenuto nell'insieme delle cose rosse.*

Ma si capisce che comunque in questo modo il coinvolgimento degli insiemi è inessenziale. Di natura completamente differente è invece il coinvolgimento degli insiemi nella frase *c*). Infatti:

non ha senso affermare che ciascuna rosa del vaso è "dodici".

Il predicato "*avere dodici elementi*" si riferisce all'intero insieme delle rose che pertanto viene ad essere il vero soggetto della frase *c*). In tale modo tale insieme viene ad assumere il carattere di sostanza individuale e diviene un nuovo soggetto (al singolare) distinto dagli elementi che lo compongono. Ora, prima di Cantor, tale oggetto non veniva considerato come ente matematico e, conseguentemente, la *c*) non risultava essere una asserzione matematica. Essa esprimeva il risultato di un esperimento (il contare) su di un oggetto e non era diversa da una frase del tipo "*la temperatura del tale corpo è di dodici gradi*". Dopo Cantor invece gli insiemi saranno visti come nuovi enti matematici e la *c*) sarà una asserzione matematica come le altre.⁶

Concludiamo questo paragrafo osservando che, in definitiva, vi sono due modi di coinvolgere una collezione di elementi in un discorso scientifico:

- Il primo si ha quando ci si riferisce a ciascun elemento della collezione stessa.
- Il secondo modo si ha quando si considera tale collezione come ente matematico a cui è possibile attribuire proprietà che non sono riconducibili ai suoi elementi.

Spesso nel primo caso si utilizza il termine "classe" e solo nel secondo caso si usa il termine "insieme" per denotare la collezione. La teoria (per meglio dire il linguaggio) delle classi è sempre esistita e non permette di dire niente di più di quanto

⁶ Ricordiamo che già c'era stato un momento dell'evoluzione della matematica in cui questa aveva allargato il proprio ambito creando nuovi enti astratti. Mi riferisco al processo di idealizzazione degli enti geometrici in cui proprietà come "*essere retto*", "*essere circolare*", "*essere quadrato*" vengono sostituite da enti astratti come "*la retta*", "*il cerchio*", "*il quadrato*".

permetta il linguaggio comune. La zoologia, la mineralogia l'hanno spesso utilizzata. Agli oggetti ed alle proprietà si sostituiscono gli elementi e le classi (estensioni di tali proprietà). Alla congiunzione logica "e" ed alla disgiunzione "o" si sostituiscono le operazioni di intersezione e di unione, alla negazione si sostituisce la complementazione, alla implicazione la relazione di inclusione.

La teoria degli insiemi nasce invece con Cantor ed il suo significato si manifesta esclusivamente nell'ambito matematico. Sue caratteristiche peculiari furono l'esame della "grandezza" degli insiemi ed il tentativo di procedere ad una fondazione di tutta la matematica. Quella che molti studenti hanno imparato nelle scuole è solo la teoria della classi, o, per meglio dire, il linguaggio delle classi.

Si fa' invece teoria degli insiemi quando (in generale nelle elementari) si parla di "cardinalità" di un insieme finito tramite il concetto di equipotenza. Ancora si utilizza la teoria degli insiemi quando si costruisce il campo dei numeri reali con il metodo delle sezioni o con un qualunque altro metodo.⁷

3. I paradossi dell'infinito

Una delle prime domande che un matematico si pone riguarda la grandezza degli enti matematici che si propone di studiare. Ed infatti la parte centrale degli studi di Cantor sugli insiemi riguar-

⁷ La distinzione che abbiamo fatto in questo paragrafo è importante anche per comprendere il significato che ha l'introduzione (eventuale) della teoria degli insiemi nelle scuole. A volte si definisce il massimo comune divisore tra i numeri n ed m come "l'elemento massimo dell'intersezione tra l'insieme dei divisori di n e l'insieme dei divisori di m ". A me sembra un modo inutilmente complicato di dare la definizione! Il linguaggio comune in questo caso è più che sufficiente e non c'è modo più semplice di definire il massimo comune divisore che affermare che è il massimo dei divisori comuni. Un altro esempio di apparente introduzione della teoria degli insiemi è quando, ad esempio, per illustrare la nozione di intersezione, si dice che la balena è un elemento di intersezione tra gli insieme degli animali acquatici e l'insieme dei mammiferi. Anche in questo caso sembra più semplice dire che la balena è un mammifero che vive nell'acqua senza disturbare l'operazione di intersezione.

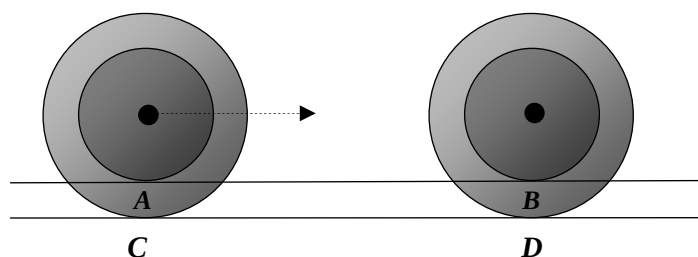
da la loro grandezza, o per meglio dire, la loro cardinalità. Siamo nel 1874 quando appaiono i primi articoli di Cantor in proposito, ma già Galileo nel 1638 nel suo *Discorsi e dimostrazioni Matematiche intorno a due nuove scienze* si era posto il problema sulla possibilità di confrontare la grandezza di due insiemi infiniti. La risposta di Galileo a questa domanda fu nettamente negativa, e questo in base ad alcune conseguenze paradossali conseguenti a tale possibilità.

Il paradosso dei quadrati perfetti. Galileo confrontò l'insieme N degli interi e l'insieme QP dei quadrati perfetti e giunse alla paradossale conclusione che tali insiemi hanno lo stesso numero di elementi. Infatti egli osservò che ad ogni intero n è possibile associare il quadrato perfetto n^2 ottenendo in tale modo la tabella

1	2	3	4	5	6	...
1	4	9	16	25	36	...

Naturalmente in tale modo ad elementi distinti corrispondono quadrati distinti ed ogni quadrato si può ottenere in questo modo (in termini moderni diremmo che la corrispondenza è biettiva) ciò prova che ci sono tanti elementi in N quanti ce ne sono in QP . Ora una tale conclusione non appariva a Galileo soltanto contraria al senso comune (paradossale) ma anche contraddittoria. Precisamente era in contraddizione con l'assioma euclideo "il tutto è maggiore della parte" che Galileo, grande ammiratore di Euclide, non si sognava di mettere in discussione.

Il Paradosso delle ruote concentriche: Questo paradosso, descritto da Galileo, risale ad Aristotele.



Consideriamo due ruote concentriche incollate una sull'altra, di raggi 0.5 ed 1 e supponiamo di far fare alla più grande un giro

completo in modo che rotolando tracci un segmento CD . Nel frattempo anche la ruota più piccola avrà fatto un giro completo ed avrà tracciato un segmento AB della stessa lunghezza. Ma questo è assurdo perché i due segmenti rappresentando lo “srotolamento” di due circonferenze di lunghezza π e 2π dovrebbero essere uno minore dell'altro. Ecco quello che dice Galileo:

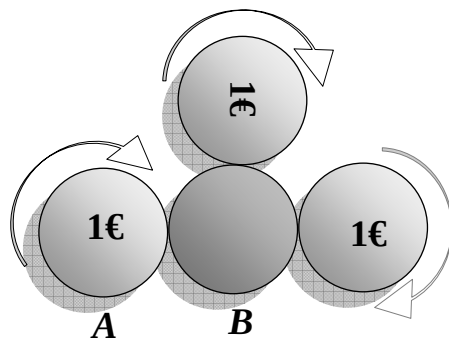
“Or come dunque può senza salti scorrere il cerchio minore una linea tanto maggiore della sua circonferenza ?”.

Comunque la conclusione a cui Galileo pervenne fu non tanto che non fosse possibile considerare l'infinito attuale (come affermato dalla tradizione aristotelica) ma che la comprensione dell'infinito attuale forse era al di fuori delle capacità dell'essere umano. Egli infatti asserisce che quando

siamo tra gli infiniti e gli indivisibili, quelli sono incomprendibili dal nostro intelletto finito per la loro grandezza, e questi per la loro piccolezza.

Dove il termine "indivisibile" si riferisce agli infinitesimi che in quel periodo, ad opera del Cavalieri, si andavano utilizzando.

Paradosso delle due monete. Un paradosso simile a quello di Aristotele si ottiene supponendo di porre su di un tavolo due monete uguali A e B in modo che si tocchino in un punto P , di tenere fissa A e di fare ruotare di un giro intero B mantenendo sempre un punto P di contatto. Ci si accorge che alla fine della rotazione B non ha percorso tutta la circonferenza di A ma solo la metà e che quindi il punto P di contatto visto come punto di A percorre una lunghezza π , visto come punto di A percorre una lunghezza 2π .



Paradosso delle serie. I paradossi che abbiamo visto fino ad ora sono legati alla nozione di “grandezza”. Altri paradossi sono legati alla nozione di “somma infinita”. Ne abbiamo già visto uno quando abbiamo parlato del paradosso di Zenone. Un altro, oggetto di discussione agli inizi del settecento, è il seguente. Consideriamo la somma infinita $+1-1+1-1+\dots$ che si ottiene alternando $+1$ e -1 e cerchiamo di capire quale è il valore di tale somma.⁸ Una possibile risposta è che sia uguale a zero. Infatti se raggruppano gli elementi di tali somme nella sequenza $(+1-1)+(+1-1)+\dots$ allora eseguendo le sottrazioni tra parentesi ed effettuando una somma di infiniti zero, si ottiene zero. Un'altra possibilità è di ottenere 1. Infatti basta isolare il primo 1 e poi raggruppare le coppie di successivi numeri ed ottenere pertanto $1+(-1+1)+(-1+1)+\dots$. Padre Guido Grandi sosteneva che la somma fosse uguale a 0.5. Infatti egli argomentava al modo seguente. Supponiamo che in una eredità di un prezioso gioiello si fosse stabilito che un anno esso dovesse essere tenuto da un fratello e l'altro anno da un altro fratello. Allora il tempo di possesso di un fratello risulterebbe uguale a quello dell'altro e quindi è uguale a 0.5.⁹ Una argomentazione diversa, ma sempre a favore del valore 0.5, era data da Leibniz che sosteneva che se si interrompe a caso il calcolo del valore della serie, allora la somma può fornire il valore 0 o 1 e questo con la stessa probabilità. Quindi il valore della somma infinita si deve collocare a metà strada tra 0 ed 1 ed è pertanto 0.5. In termini attuali diremmo che il valore di aspettazione della sequenza è 0.5.

4. Cantor, l'infinito e la dottrina Cristiana

Cantor, al contrario di Galileo, non esitò a gettare via l'assioma di Euclide accettando tranquillamente la possibilità che esistano insiemi che hanno tanti elementi quanto una loro parte propria.¹⁰

⁸ Se ci si riferisce alla definizione attuale di somma di una serie, tale serie risulta non convergente.

⁹ Il fatto che fosse possibile fare passare il valore della serie da zero ad 1 oppure a 0.5 veniva visto da padre Grandi non come una contraddizione ma come prova che Dio, nella sua infinità, potesse creare dal nulla il tutto.

¹⁰ Questa proprietà diventerà proprio un modo di definire gli insiemi infiniti. Quindi si potrebbe dire che un insieme è infinito se verifica il paradosso di Galileo.

E' da notare che le motivazioni filosofiche di Cantor erano strettamente intrecciate a quelle religiose. La cosa non deve sorprendere poiché le motivazioni religiose stanno alla base dei primi esempi di accettazione dell'infinito attuale nella cultura occidentale. Era infatti convincimento comune nel medioevo che Dio fosse infinito.

Appunto perché è uno, Egli non rientra né in una misura né in un numero. Così Egli non incontra il confine né in altrui né in se stesso, che, in tal caso, Egli cadrebbe già nella dualità. (Plotino, Enneadi).

Era semmai oggetto di discussione se Dio potesse concepire entità infinite poiché in tale caso si dovrebbero accettare entità infinite diverse da Dio. A tale proposito ad esempio l'opinione di San Tommaso era che l'unico infinito fosse Dio.

Quindi, come Dio, nonostante abbia potenza infinita, tuttavia non può creare qualcosa di increato (il che sarebbe far coesistere cose contraddittorie), così non può creare cosa alcuna che sia assolutamente infinita. (S. Tommaso, Summa Teologica).

In altre parole, se si vede che il concetto di infinito attuale è contraddittorio, Dio non può pensarlo perché Dio, in un certo senso, rispetta la logica. Invece l'opinione di Sant'Agostino era non solo che Dio fosse infinito ma anche che potesse avere come oggetto del suo pensiero "il tutto del numero" cioè l'intera totalità degli interi.

Riguardo poi all'altra loro teoria che neanche con la scienza di Dio può essere rappresentato l'infinito, rimane loro che o sino affermare, immergendosi nell'abisso profondo della irreligiosità, che Dio non conosce il tutto del numero . . . Non lo potrebbe dire neanche il più insensato . . . Che razza di omucci siamo noi che pretendiamo di porre limiti alla sua scienza?. (Agostino, La città di Dio)

Ad esempio in una lettera del 1890 a padre Thomas Esser egli scrive tra l'altro:

Viene da me offerta alla filosofia cristiana per la prima volta la vera dottrina dell'infinito nei suoi principi. So con piena si-

curezza e determinazione che essa accoglierà questa dottrina: è soltanto da vedere se ciò accadrà già adesso o soltanto dopo la mia morte.

Da notare che Cantor distingue tre tipi di infinito. Il primo, legato all'idea di Dio, il secondo di natura fisica (il tempo, lo spazio), il terzo di natura matematica.

L'infinito attuale si presenta in tre contesti: il primo è quello in cui si presenta nella forma più completa, in un essere completamente indipendente trascendente questo mondo, "in Deo", ed è questo che io chiamo l'Infinito Assoluto o semplicemente l'Assoluto; il secondo quando si presenta nel mondo contingente, nel creato; il terzo è quando la mente lo afferra "in abstracto", come grandezza matematica, numero o tipo d'ordine. Voglio sottolineare chiaramente la differenza tra l'Assoluto e quello che io chiamo il Transfinito, cioè l'infinito attuale degli ultimi due tipi, perché si tratta di oggetti evidentemente limitati, suscettibili di accrescimento, e quindi collegati al finito.

5. Confrontare le grandezze degli insiemi.

La possibilità di confrontare le grandezze di insiemi non necessariamente finiti è alla base della teoria degli insiemi.

Definizione 5.1. Diciamo che due insiemi A ed B sono equipotenti e scriviamo $A \equiv B$ se esiste una corrispondenza biettiva $f: A \rightarrow B$ di A in B . Diciamo che la potenza di A è minore o uguale della potenza di B e scriviamo $A \leq B$ se esiste una corrispondenza iniettiva di A in B .

In altri termini A ed B sono equipotenti se esiste un procedimento che permette di ottenere, a partire da un elemento di A , uno ed un solo elemento di B in modo che

- ad elementi distinti di A corrispondano elementi distinti di B
- ogni elemento di B si può ottenere in questo modo.

Nel caso della corrispondenza proposta da Galileo tra N e l'insieme QP dei quadrati perfetti, tale corrispondenza era ottenuta tramite l'elevazione al quadrato. Ovviamente nel caso di insiemi finiti si ha che due insiemi sono equipotenti se e solo se

hanno lo stesso numero di elementi nel senso intuitivo che diamo a questa espressione.

Teorema 5.2. Comunque si considerino gli insiemi A , B e C :

1. $A \equiv A$.
2. $A \equiv B$ implica $B \equiv A$.
3. $A \equiv B$ e $B \equiv C$ implica $A \equiv C$.

Pertanto, in un certo senso, la relazione di equipotenza è una relazione di equivalenza.¹¹

Dim. 1. Sia A un qualunque insieme, allora l'applicazione identica $i : A \rightarrow A$ è una funzione biettiva di A in A . Quindi A è equipotente ad A .

2. Sia A equipotente a B , allora esiste una corrispondenza biettiva $f : A \rightarrow B$. La sua funzione inversa sarà allora una corrispondenza biettiva di B in A e questo prova che B è equipotente ad A .

3. Supponiamo che A abbia potenza uguale a B e che B abbia potenza uguale a C , allora esistono due funzioni biettive $f : A \rightarrow B$ e $g : B \rightarrow C$. È evidente che fg è una funzione biettiva di A in C e quindi che A e C sono equipotenti. $\square$

Teorema 5.3. La relazione $\leq$ “avere meno potenza di” è una relazione di pre-ordine, cioè comunque si scelgano gli insiemi A , B e C :

- i) $A \leq A$ (riflessiva)
 - ii) Se $A \leq B$ e $B \leq C$ implica $A \leq C$ (transitiva).
- L'equivalenza associata a tale pre-ordine è l'equipotenza, cioè
- iii) $A \leq B$ e $B \leq A$ implica $A \equiv B$ (Teorema di Cantor-Bernstein).

Dim. Per provare i) e ii) si procede allo stesso modo che nel Teorema 5.2. Per dimostrare iii) si veda il Paragrafo 12.

¹¹ In realtà non è completamente corretto parlare di relazione di equivalenza. Infatti l'insieme in cui tale relazione dovrebbe essere considerata dovrebbe essere la classe di tutti gli insiemi. Purtroppo, come vedremo nel prossimo capitolo, tale classe crea molti problemi. Analoga considerazione deve essere fatta per la relazione “avere meno potenza di” che viene trattata nel teorema successivo. Insomma le cose sono sempre più complicate di come si pensa.

Proposizione 5.4. Sia A un insieme e $\equiv$ una relazione di equivalenza in A , allora il quoziente $A/\equiv$ ha potenza minore o uguale ad A .

Dim. Sia $f : A/\equiv \rightarrow A$ una funzione che associ ad ogni classe di equivalenza z un elemento $f(z) \in z$.¹² La funzione f è iniettiva perché se $f(z) = f(z') = c$ allora le due classi z e z' hanno l'elemento c in comune e quindi coincidono. $\square$

Proposizione 5.5. Se esiste una funzione suriettiva $f : A \rightarrow B$ di A in B allora B ha potenza minore o uguale ad A . Infatti se $\equiv$ è il nucleo¹³ di f , allora B è equipotente a $A/\equiv$.

Dim. Sia $f : A \rightarrow B$ una funzione suriettiva, allora il suo nucleo $\equiv$ ripartisce A in classi di equivalenza. La funzione $g : A/\equiv \rightarrow B$ definita ponendo $g([x]) = f(x)$ è ben definita perché il suo valore non dipende dall'elemento rappresentativo in $[x]$. Inoltre, poiché $g([x]) = g([y]) \Rightarrow f(x) = f(y) \Rightarrow x \equiv y \Rightarrow [x] = [y]$, la funzione g è iniettiva. Pertanto $A/\equiv$ è equipotente a B . $\square$

Il paradosso di Galileo di un insieme equipotente ad una sua parte propria sfrutta il fatto che i numeri naturali sono un insieme infinito. Come abbiamo già osservato in una nota, questo fenomeno può essere utilizzato proprio per dare una definizione di insieme infinito.

Definizione 5.6. Un insieme equipotente ad una propria parte propria viene detto *infinito*, altrimenti viene detto *finito*.

Proposizione 5.7. L'insieme N dei numeri naturali è infinito. Inoltre un insieme è infinito se e solo se ha potenza maggiore o uguale ad N .

Dim. Che N sia infinito deriva dall'osservazione di Galileo per cui N è equipotente all'insieme QP dei quadrati perfetti che costituiscono una parte propria di N . D'altra parte ogni terna di Peano

¹² Che tale funzione esista equivale all'assioma della scelta di cui parleremo più in avanti.

¹³ Il nucleo di una funzione f è l'insieme delle coppie (x, y) tali che $f(x) = f(y)$, in proposito si veda in Appendice.

è infinita perché la funzione successore è iniettiva per definizione e non è suriettiva in quanto il primo elemento non è successore di nessun altro elemento.

Sia S un insieme infinito, allora abbiamo già visto che all'interno di N è possibile definire una terna di Peano. Pertanto, essendo tutte le terne di Peano isomorfe e quindi equipotenti, S ha potenza maggiore di N . Viceversa se S ha potenza maggiore o uguale al numerabile allora esiste una funzione iniettiva $f : N \rightarrow S$. Consideriamo la funzione $g : S \rightarrow S$ definita ponendo $g(x) = x$ se $x \notin f(N)$ e $g(f(n)) = f(n+1)$: in altri termini tale funzione sposta di un passo gli elementi della successione $f(n)$ e lascia immutati gli altri elementi di S . Ne segue che, poiché g non può assumere il valore $f(1)$, g è una funzione biettiva di S nella sua parte propria $S - f(1)$. Ciò prova che S è infinito. $\square$

Problema. Dimostrare che a Salerno esistono due persone che hanno esattamente lo stesso numero di capelli. E' necessario sapere che:

- E' stato osservato che un uomo non ha mai più di 130.000 capelli.
- si sa che a Salerno vivono più di 200.000 persone.

6. Insiemi numerabili

N è il primo esempio di insieme infinito che abbiamo incontrato. Confrontiamo allora la grandezza degli insiemi con quella di N .

Definizione 6.1. Gli insiemi equipotenti ad N vengono detti *numerabili*, gli insiemi che hanno potenza minore od uguale ad N vengono chiamati *enumerabili*.¹⁴

Pertanto un insieme A è numerabile se esiste una funzione $f : A \rightarrow N$ biettiva. A è enumerabile se esiste una funzione $f : A \rightarrow N$ iniettiva. Se si tiene conto della Proposizione 5.5 abbiamo la seguente ovvia proposizione.

¹⁴ In realtà il concetto di insieme enumerabile non sembra essere usato nei libri di matematica italiani. Nei libri in lingua inglese invece viene trattato sotto il nome di *countable set*. Invece in informatica teorica si parla di *effettivamente enumerabile* (in inglese *effectively enumerable*) per indicare un insieme che sia enumerato tramite una funzione per la quale esiste un opportuno programma capace di computarla.

Proposizione 6.2. Le seguenti asserzioni sono equivalenti:

- a) A è enumerabile
- b) A è vuoto, finito o numerabile
- c) A è vuoto oppure esiste una funzione $f: N \rightarrow A$ suriettiva (che viene chiamata *funzione enumerante*).

In termini intuitivi un insieme non vuoto è enumerabile se si possono "numerare" tutti gli elementi di A uno dopo l'altro in una successione dicendo che

$f(1)$ è il primo elemento
 $f(2)$ è il secondo elemento

...

ed in tale numerazione nessun elemento di A deve sfuggire (f è suriettiva). Se in tale enumerazione non esistono ripetizioni, essendo la f anche iniettiva, A risulta numerabile.

Proposizione 6.3. L'unione di due insiemi enumerabili è un insieme enumerabile. L'unione di un insieme finito ed uno numerabile è un insieme numerabile. L'unione di due insiemi numerabili è un insieme numerabile.¹⁵

Dim. Siano A e B due insiemi enumerabili e siano $f: N \rightarrow A$ e $g: N \rightarrow B$ due funzioni enumeranti A e B . Intuitivamente il processo di enumerazione di $A \cup B$ consiste nell'alternare la numerazione di A e di B , cioè nel considerare la seguente successione

$f(1), g(1), f(2), g(2), \dots$

Più precisamente, possiamo considerare la funzione $h: N \rightarrow A \cup B$ definita ponendo

$$h(n) = \begin{cases} f((n+1)/2) & \text{se } n \text{ è dispari} \\ g(n/2) & \text{se } n \text{ è pari.} \end{cases}$$

Tale funzione è suriettiva e quindi h rappresenta una funzione enumerante $A \cup B$.

Sia $A = \{a_1, \dots, a_p\}$ un insieme finito e B un insieme numerabile. Sia $f: N \rightarrow B$ una funzione biettiva enumerante B , allora una numerazione di $A \cup B$ si ottiene al modo seguente:

¹⁵ Alla fine di questo capitolo tale teorema viene illustrato tramite il racconto degli "alberghi di Hilbert".

$$a_1, a_2, \dots, a_p, f(1), f(2), \dots$$

In altri termini si cominciano ad elencare tutti gli elementi di A e poi si elencano quelli di B . Volendo formalizzare un tale modo di procedere definiamo la funzione $h : N \rightarrow A \cup B$ ponendo

$$\begin{aligned} h(1) &= a_1, h(2) = a_2, \dots, h(p) = a_p, \\ h(p+1) &= f(1), \dots, h(p+m) = f(m), \dots \end{aligned}$$

E' evidente che h è una funzione biettiva enumerante $A \cup B$.

Siano A e B numerabili, allora nel caso in cui A e B siano disgiunti la funzione h ora definita è biettiva. Invece nel caso $A \cap B \neq \emptyset$ h non è iniettiva in quanto un elemento dell'intersezione viene enumerato due volte. Intuitivamente basta che nella enumerazione

$$f(1), g(1), f(2), g(2), \dots$$

gli elementi già comparsi non siano ripetuti.¹⁶ □

Esempio. Sia $A = \{5, 3, 7\}$ e sia QP l'insieme infinito dei quadrati perfetti. Allora possiamo enumerare al modo seguente gli elementi di $A \cup QP$:

$$5, 3, 7, 1, 4, 9, \dots$$

Più precisamente sappiamo che la funzione $f(n) = n^2$ è una biiezione di N in QP , cioè è una funzione enumerante QP . Possiamo allora definire una funzione $h : N \rightarrow A \cup QP$ enumerante $A \cup QP$ ponendo:

$$\begin{aligned} h(1) &= 5 \\ h(2) &= 3 \\ h(3) &= 7 \\ h(4) &= f(1) = 1 \\ h(5) &= f(2) = 4 \\ &\dots \\ h(n) &= f(n-3). \end{aligned}$$

Proposizione 6.4. L'insieme Z degli interi relativi è numerabile.¹⁷

¹⁶ Un modo più rigoroso ma meno costruttivo di procedere è quello di osservare che, essendo h suriettiva $A \cup B$ ha potenza minore o uguale a quella di N . Poiché è evidente che $A \cup B$ ha potenza maggiore o uguale a quella di N , $A \cup B$ è equipotente ad N . Tale modo di ragionare non fornisce però concretamente la funzione enumerante.

Dim. L'insieme Z è unione dell'insieme degli interi positivi Z^+ (compreso lo zero) e dell'insieme degli interi negativi Z^- . Poiché entrambi questi due insiemi sono numerabili, Z è numerabile. Più precisamente una enumerazione degli elementi di Z si ottiene al modo seguente

0, -1, +1, -2, +2, 3, -3 ...

cioè tramite la funzione $h : N \rightarrow Z$ definita ponendo $h(n) = -n/2$ se n è pari e $h(n) = (n-1)/2$ se n è dispari. $\square$

7. Tentare di superare il numerabile

Ed io sono più infinito di te !

No ! più infinito di me non esiste niente.

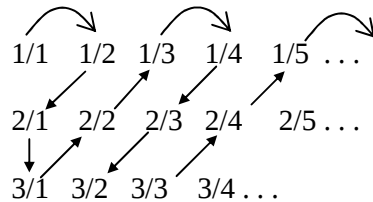
(Da una disputa tra Il Dio del Nord ed Il Dio del Sud prima dello scontro di civiltà del 2013)

Naturalmente si pone il problema se sia possibile trovare insiemi che siano più grandi del numerabile. La prima cosa che viene in mente naturalmente è quella di confrontare N con l'insieme Q dei razionali. Inaspettatamente Cantor provò che:

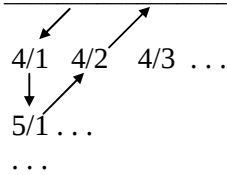
anche Q , che appariva tanto più grande di N , aveva tanti elementi quanti ne aveva N !

Proposizione 7.1. L'insieme Q dei numeri razionali è numerabile.

Dim. Cominciamo col provare che l'insieme dei razionali positivi Q^+ è numerabile. Una dimostrazione intuitiva di tale numerabilità si ottiene disponendo i numeri razionali nella seguente matrice infinita



¹⁷ Per trovare altre dimostrazioni di questa proposizione e delle successive si consiglia di vedere il paragrafo 3 del Capitolo 5 dove sono presenti anche altri esempi di insiemi numerabili



in cui il numeratore rappresenta la riga di appartenenza ed il denominatore la posizione all'interno della riga di un razionale. Ovviamente un tale matrice un razionale compare più volte. Il problema è trovare una strategia che consenta di *leggere* i razionali in questa tabella uno dopo l'altro in modo che:

- nessun razionale sia escluso (la corrispondenza deve essere suriettiva)
- nessun razionale sia letto due volte (la corrispondenza deve essere iniettiva).

E' possibile seguire la strategia indicata dalle frecce avendo cura di *saltare* i numeri razionali che già si sono incontrati. Per precisare una tale strategia osserviamo che si è cominciato a "contare" gli elementi dell'insieme $Q_2 = \{1/1\}$, poi si sono contati gli elementi dell'insieme $Q_3 = \{1/2, 2/1\}$, poi gli elementi dell'insieme $Q_4 = \{3/1, 2/2, 1/3\}$... Ciascuno di tali insiemi si può caratterizzare dal fatto che la somma del numeratore e del denominatore è rispettivamente 2, 3, Indichiamo con Q_n l'insieme dei razionali del tipo p/q con p e q interi positivi primi tra loro e $p+q = n$. Pertanto,

$$\begin{aligned} Q_2 &= \{1/1\} = \{1\}, \\ Q_3 &= \{1/2, 2/1\}, \\ Q_4 &= \{1/3, 3/1\}, \\ Q_5 &= \{1/4, 4/1, 2/3, 3/2\} \\ &\dots \end{aligned}$$

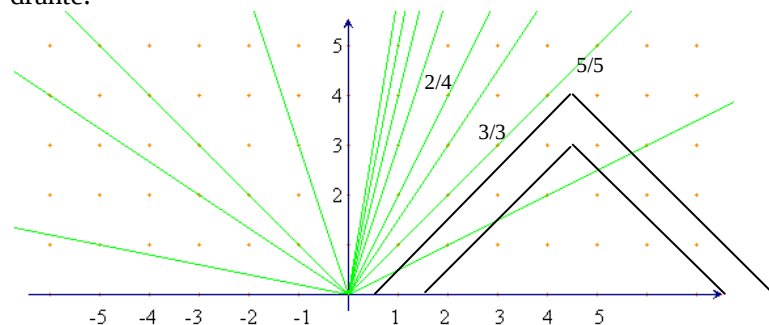
E' evidente che ciascun Q_n è finito e che $Q^+ = \bigcup_{n \in \mathbb{N} - \{1\}} Q_n$. Una numerazione degli elementi di Q^+ si ottiene pertanto enumerando prima gli elementi di Q_2 , poi quelli di Q_3 e così via. In tale modo si ottiene la numerazione:

$$0, 1, 1/2, 2/1, 1/3, 3/1, 1/4, 4/1, 2/3, 3/2, \dots$$

Una volta che si è visto che Q^+ è numerabile risulta evidente che anche Q^- è numerabile in quanto equipotente a Q^+ . Quindi $Q =$

$Q^+ \cup Q^- \cup \{0\}$ è numerabile in quanto unione di un numero finito di insiemi finiti o numerabili. $\square$

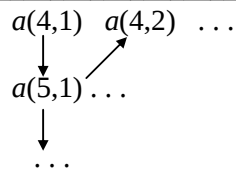
Possiamo dare un carattere più geometrico alla dimostrazione ora esposta. Basta associare ad ogni punto del piano cartesiano le cui coordinate sono (n,m) , con $m \neq 0$, il corrispondente numero razionale n/m . Ovviamente tale corrispondenza non è iniettiva ma punti a cui corrispondono lo stesso razionale si dispongono su di una retta passante per l'origine. Questo è un modo con cui a volte si introducono i razionali nelle scuole medie. I razionali positivi si collocano nel primo quadrante, quelli negativi nel quarto quadrante.



La costruzione fatta per enumerare Q^+ sostanzialmente abbiamo considerato una successione di triangoli rettangoli con un vertice nell'origine e gli altri due vertici nei punti $(n,0)$ e $(0,n)$. Per enumerare i razionali sia positivi che negativi è possibile considerare la successione di triangoli di vertici $(-n,0)$, $(0,n)$, $(n,0)$. Avremmo anche potuto considerare successioni di altre figure geometriche limitate, ad esempio una successione di opportuni cerchi di centro l'origine.

Proposizione 7.2. Il prodotto cartesiano di un numero finito di insiemi numerabili è un insieme numerabile.¹⁸

¹⁸Per enumerare $N \times N$ esiste anche una funzione di cui si ha una espressione analitica. Infatti basta considerare la funzione $h : N \times N \rightarrow N$ definita ponendo $h(i,j) = 2^i \cdot (2 \cdot j + 1)$. Tale funzione è biettiva perché un numero n può essere scomposto in uno ed un solo modo come prodotto del tipo $2^n \cdot (2 \cdot m + 1)$. Tale scomposizione si ottiene calcolando la massima potenza 2^i di 2 che divide n . Allora $n = 2^i \cdot d$ con d dispari. Posto $j = (d-1)/2$ risulta che $n = 2^i \cdot (2 \cdot j + 1) = h(i,j)$.



8. La potenza del continuo

Sembra di essere ad un punto morto: abbiamo confrontato l'insieme QP dei quadrati perfetti con N ed abbiamo visto che hanno lo stesso numero di elementi, poi abbiamo esaminato Z , Q^+ e Q ed ancora abbiamo visto che sono numerabili. Si pone allora il seguente interrogativo:

e se tutti gli insiemi infiniti fossero numerabili?

In tale caso una teoria che si occupa delle grandezze degli insiemi si ridurrebbe a ben poca cosa perché avremmo da un lato gli insiemi finiti, di diverse grandezze, e dall'altra tutti i rimanenti insiemi tutti della stessa grandezza infinita. D'altra parte è questa l'opinione di Galileo che afferma.

Io non veggio che ad altra decisione si possa pervenire, che a dire, infiniti essere tutti i numeri, infiniti i quadrati, infinite le lor radici, ne' la moltitudine dei quadrati essere minore di quella di tutti i numeri, ne' questa maggiore di quella, ed in ultima conclusione, gli attributi di eguale, maggiore e minore non aver luogo negli infiniti, ma solo nelle quantità terminate.

Ma il merito di Cantor non fu solo di avere avuto il coraggio di negare validità all'assioma euclideo "il tutto è maggiore della parte" e di accettare l'esistenza dell'infinito attuale. Egli riuscì anche a trovare due insiemi infiniti di grandezza diversa; precisamente riuscì a provare che l'insieme R dei numeri reali ha potenza maggiore di quella di N e che quindi esistono almeno due ordini di infinito. Cominciamo con il dimostrare la non numerabilità dell'intervallo aperto $(0,1)$. Il metodo usato in questa dimostrazione viene a volta chiamato *metodo di diagonalizzazione di Cantor* ed è stata proposto da Cantor nel 1891.¹⁹

¹⁹ In realtà la prima dimostrazione, dovuta sempre a Cantor, risale al 1873 ed è di tipo topologico. Supponiamo per assurdo che $f: N \rightarrow (0,1)$

Teorema 8.1. L'insieme $(0,1)$ non è numerabile e quindi è più grande di N .

Dim. Supponiamo per assurdo che esista una funzione biettiva $f: N \rightarrow (0,1)$. Come è noto ogni numero $f(n)$ della successione $f(1), f(2), \dots$ può essere rappresentato con una espansione decimale e ciò consente di costruire una tabella del tipo

$$\begin{aligned} f(1) &= 0. \mathbf{r(1,1)} r(1,2) r(1,3) \dots \\ f(2) &= 0. r(2,1) \mathbf{r(2,2)} r(2,3) \dots \\ f(3) &= 0. r(3,1) r(3,2) \mathbf{r(3,3)} \dots \\ &\dots \\ f(n) &= 0. r(n,1) r(n,2) r(n,3) \dots \mathbf{r(n,n)} \dots \\ &\dots \end{aligned}$$

dove $r(n,i)$ rappresenta la cifra di posto i -esimo della espansione decimale di $f(n)$. Si noti che, stante il significato del periodo 9 nella rappresentazione decimale di un numero, uno stesso numero può essere indicato con espansioni diverse, ad esempio $0.2999\dots$ coincide con il numero $0.3000\dots$. Supporremo di rappresentare i numeri sempre senza il periodo nove in modo che espansioni diverse rappresentino sempre numeri diversi.²⁰ Ci

sia una funzione biettiva. Allora possiamo definire una successione $(I_n)_{n \in N}$ di intervalli chiusi contenuti in $(0,1)$ tale che

- I_1 non contiene $f(1)$
- $I_2 \subseteq I_1$ e non contiene $f(2)$
- ...
- $I_{n+1} \subseteq I_n$ e non contiene $f(n+1)$
- ...

Poiché $(I_n)_{n \in N}$ è una successione decrescente di compatti non vuoti, la sua intersezione è non vuota. Detto r un numero reale in $\bigcap_{n \in N} I_n$ è evidente che r è un elemento di $(0,1)$ diverso da ogni $f(n)$ in contrasto con l'ipotesi di suriettività per f . L'assurdo a cui siamo pervenuti prova che $(0,1)$ non è numerabile. Il fatto che tale dimostrazione sia di natura squisitamente topologica non deve sorprendere se si pensa che Cantor ha sviluppato le proprie idee nell'ambito delle sue ricerche di analisi matematica. Da un punto di vista didattico il pregio di questa dimostrazione è che non presuppone il teorema di rappresentazione dei numeri reali in forma decimale. Il difetto è che utilizza la nozione di compattezza che non è certo semplice.

²⁰ Molti miei studenti agli esami dicono una frase del tipo "escludiamo i numeri di periodo nove". La dimostrazione funziona lo stesso ma in

proponiamo ora di costruire un numero reale $r = 0.c_1c_2c_3\dots$ appartenente a $(0,1)$ che sia diverso da tutti gli elementi della successione $f(1), f(2), \dots$. A tale scopo basta prendere c_1 diverso da $r(1,1)$, c_2 diverso da $r(2,2)$ e, più in generale, c_n diverso da $r(n,n)$. Ad esempio possiamo porre

$$c_n = \begin{cases} 1 & \text{se } r(n,n) \neq 1 \\ 2 & \text{se } r(n,n) = 1. \end{cases}$$

Allora, per come è stato costruito, il numero r non può coincidere con nessun numero $f(n)$, in contrasto con l'ipotesi per cui f è suriettiva.

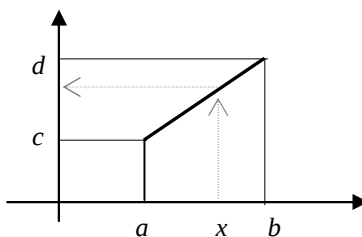
Problema. Dire perché il numero r definito all'interno della dimostrazione della Proposizione 7.1 non coincide con $f(5)$.

In tale proposizione avremmo potuto considerare anche intervalli diversi da $(0,1)$. D'altra parte la seguente proposizione mostra che tutti gli intervalli aperti sono equipotenti tra loro.

Proposizione 8.2. Tutti gli intervalli chiusi sono equipotenti tra loro. Tutti gli intervalli aperti sono equipotenti tra loro.

Dim. Dati gli intervalli chiusi $[a,b]$ e $[c,d]$, consideriamo la retta passante per il punto di coordinate (a,c) e (b,d) . Tale retta è il grafico della funzione

$$f(x) = \frac{d-c}{b-a}(x-a) + c$$



E' immediato che tale funzione è una corrispondenza biettiva tra $[a,b]$ e $[c,d]$.

La stessa funzione f può essere vista come una corrispondenza biettiva tra gli intervalli aperti (a,b) e (c,d) e questo prova che anche due intervalli aperti sono equipotenti tra loro. $\square$

questo modo quello che si dimostra è che l'insieme dei numeri "che non hanno periodo nove" non è numerabile. Naturalmente da ciò si ricava che, a maggior ragione, $[0,1]$ non è numerabile e quindi il tutto funziona lo stesso. Tuttavia questo modo di dire nasconde la difficoltà di distinguere un numero reale (che è una classe) dalla sua rappresentazione.

Problema. Trovare l'espressione analitica di una funzione che esprima l'equipotenza tra l'intervallo $[1,2]$ e l'intervallo $[3,7]$.

Dobbiamo ora confrontare la grandezza di un intervallo aperto con quella di un intervallo chiuso. A tale scopo abbiamo bisogno della seguente proposizione.

Proposizione 8.3. Se S è infinito ed A è finito o numerabile allora S ed $S \cup A$ sono equipotenti.²¹

Dim. Abbiamo già dimostrato un teorema analogo nel caso in cui S sia numerabile. Con un piccolo adattamento la stessa dimostrazione può essere estesa al caso di un qualunque insieme infinito S . Infatti essendo S infinito la potenza di S è maggiore o uguale al numerabile. Pertanto esiste una funzione iniettiva $f : N \rightarrow S$. Sia $A = \{a_1, \dots, a_n\}$, allora possiamo definire una corrispondenza $h : S \cup A \rightarrow S$ ponendo

$$h(a_1) = f(1)$$

$$h(a_2) = f(2)$$

...

$$h(a_n) = f(n)$$

$$h(f(1)) = f(1+n)$$

...

$$h(f(i)) = f(i+n)$$

...

$$h(x) = x \quad \text{se } x \notin f(N) \text{ e } x \notin \{a_1, \dots, a_n\}.$$

In altri termini si “*spostano in avanti di n passi*” gli elementi della successione $f(n)$ e nei posti che si sono liberati “*si collocano*” gli elementi di A . E' immediato che h è iniettiva e suriettiva e che quindi $S \cup A$ è equipotente ad S .

Nel caso A numerabile, supponiamo che $g : N \rightarrow A$ sia una funzione enumerante A . Allora possiamo definire h al modo seguente:

$$h(g(n)) = f(2 \cdot n)$$

$$h(f(n)) = f(2 \cdot n + 1)$$

²¹ Possiamo visualizzare la dimostrazione riferendoci agli “alberghi di Hilbert” esposti alla fine del capitolo. Infatti se S è infinito allora ha abbastanza spazio da poter contenere un albergo di Hilbert. Ma in tale albergo anche se pieno, possiamo sempre inserire gli elementi di un insieme finito o numerabile A .

$h(x) = x$ se x non appartiene né ad $f(N)$ né a $g(N)$.

In altri termini:

- si sposta ciascun elemento $f(n)$ in $f(2 \cdot n + 1)$ e si liberano gli elementi di posto pari della successione $f(n)$.
- si utilizzano i posti liberati per inserire gli elementi della successione $g(n)$, cioè gli elementi di A
- si lasciano immutati tutti gli altri elementi. $\square$

Esercizio. Dare un esempio di funzione biettiva tra l'insieme R dei numeri reali e l'insieme $\{a, b, c\} \cup R$.

Corollario 8.4. Tutti gli intervalli sono equipotenti tra loro.

Dim. Ogni intervallo aperto (a, b) è infinito, pertanto poiché $[a, b] = (a, b) \cup \{a, b\}$, per la proposizione precedente (a, b) è equipotente a $[a, b]$. $\square$

Esempio: Supponiamo di volere definire esplicitamente una funzione biettiva f da $[0, 1]$ su $(0, 1)$. Allora basta

- fissare una qualunque funzione iniettiva $h : N \rightarrow (0, 1)$ ad esempio la funzione $h(n) = 1/(n+1)$
- definire $f : [0, 1] \rightarrow (0, 1)$ ponendo

$$\begin{aligned} f(0) &= h(1) = 1/2 \\ f(1) &= h(2) = 1/3 \\ f(1/(n+1)) &= h(n+2) = 1/(n+3) \\ f(x) &= x \quad \text{se } x \neq 1/n. \end{aligned}$$

Abbiamo visto che gli intervalli, che possono essere interpretati come segmenti, sono tutti equipotenti tra di loro ed hanno tutti potenza maggiore del numerabile. Se vogliamo cercare insiemi che hanno potenza più grande possiamo, per esempio, prendere in esame l'intero insieme dei numeri reali o, se si vuole, l'insieme dei punti di una retta. Infatti il fatto che tale insieme non sia limitato potrebbe far pensare che contenga più elementi di un intervallo. Vale però la seguente proposizione.

Proposizione 8.5. L'insieme R dei numeri reali è equipotente all'intervallo aperto $(-1, 1)$ e quindi ad un qualunque intervallo.

Dim. La funzione $f(x) = x/(x^2-1)$ porta l'intervallo $(-1,+1)$ nell'insieme R dei numeri reali in modo biiettivo (studiare il grafico della funzione per rendersene conto). $\square$

Definizione 8.6. Diciamo che un insieme ha la *potenza del continuo* se è equipotente ad R .

Abbiamo visto che tutti gli intervalli, che siano aperti o chiusi, hanno la potenza del continuo.

Problema. Dimostrare che tutte le circonferenze hanno la potenza del continuo.

Il paradosso delle ruote concentriche. Torniamo al paradosso delle ruote concentriche in cui sono coinvolte le linee r ed s su cui scorrono i punti A e C delle due circonferenze. Una possibile soluzione può essere trovata proprio con la nozione di equipotenza. Infatti è vero che possiamo considerare la corrispondenza che associa ad ogni punto della circonferenza della ruota piccola il punto del segmento AB che viene toccato in un dato istante. Tale corrispondenza è iniettiva perché due punti diversi della ruota toccheranno due punti diversi del segmento. E' anche suriettiva in quanto ogni punto del segmento sarà toccato in un dato istante. Tuttavia l'esistenza di tale corrispondenza mostra che la circonferenza piccola è equipotente al segmento AB e non che ha la sua stessa lunghezza di tale segmento.

In modo analogo possiamo risolvere il paradosso delle due monete. Infatti non deve stupire il fatto che una circonferenza sia equipotente ad una semicirconferenza.

9. Superare la potenza del continuo

Abbiamo visto che R è un insieme la cui cardinalità è più grande di quella di N . Si pone allora il problema di trovare insiemi di cardinalità maggiore di quella di R . Per prima cosa tentiamo con l'insieme delle parti di N .

Proposizione 9.1. L'insieme $\mathcal{P}(N)$ delle parti di N è equipotente all'insieme $\{0,1\}^N$ che a sua volta è equipotente all'intervallo $(0,1)$. Pertanto $\mathcal{P}(N)$ ha la potenza del continuo.

Dim. Per provare che $\{0,1\}^N$ è equipotente a $(0,1)$ osserviamo che possiamo associare ad ogni elemento $c : N \rightarrow \{0,1\}$ di $\{0,1\}^N$ il numero reale $0.c(1)c(2)\dots$. Una tale corrispondenza è iniettiva (anche se non è suriettiva) e dimostra che la potenza di $\{0,1\}^N$ è minore della potenza di $(0,1)$. Viceversa, consideriamo la corrispondenza che ad ogni numero reale $0.c(1)c(2)\dots$, che supponiamo scritto in base due e senza il periodo 1, associa la successione $c : N \rightarrow \{0,1\}$ in cui $c(n)$ è l' n -esima cifra binaria. Una tale corrispondenza è iniettiva e ciò prova che la potenza di $(0,1)$ è minore della potenza di $\{0,1\}^N$. In conclusione $\{0,1\}^N$ è equipotente a $(0,1)$ e quindi ha la potenza del continuo.

Per provare che $\mathcal{P}(N)$ è equipotente a $\{0,1\}^N$, è sufficiente considerare la funzione $H : \mathcal{P}(N) \rightarrow \{0,1\}^N$ che associa ad ogni sottoinsieme X di N la sua funzione caratteristica $c_X : N \rightarrow \{0,1\}$ definita dal porre:

$$c_X(x) = \begin{cases} 1 & \text{se } x \in X \\ 0 & \text{altrimenti.} \end{cases} \quad \square$$

Proviamo ora a vedere se l'insieme delle funzioni di N in N è più grande di R . La seguente proposizione mostra che la risposta è negativa.

Proposizione 9.2. L'insieme N^N delle funzioni di N in N ha la potenza del continuo.

Dim. Sappiamo che $N \times N$ è numerabile e quindi che $\mathcal{P}(N \times N)$ ha la potenza del continuo. D'altra parte $N^N \subseteq \mathcal{P}(N \times N)$ in quanto ogni funzione di N in N , in quanto relazione binaria, è un sottoinsieme di $N \times N$. Ciò prova che la potenza di N^N è minore o uguale a quella del continuo. D'altra parte $\{0,1\}^N$ è un sottoinsieme di N^N e quindi N^N ha potenza maggiore o uguale al continuo. $\square$

Tutti gli insiemi di cui abbiamo parlato fino ad ora sono del tipo "lineare" ed abbiamo visto che tutti hanno la potenza del continuo. Allora nella nostra ricerca di insiemi più grandi appare naturale rivolgerci ad insiemi di dimensione maggiore, ad esempio ai quadrati oppure ai cubi. Ebbene si scoprì che anche l'insieme dei punti interni ad un quadrato ha la potenza del continuo !

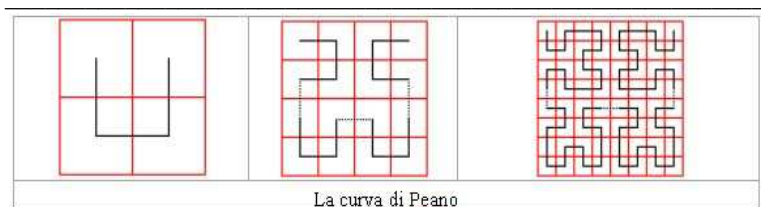
Teorema 9.3. L'insieme dei punti interni ad un quadrato è equipotente all'insieme dei punti interni ad un suo lato. Pertanto tale insieme ha la potenza del continuo.

Dim. In termini analitici dobbiamo provare che $(0,1)^2$ è equipotente a $(0,1)$. Supponiamo di rappresentare i numeri reali nella forma dell'espansione decimale. Inoltre, per ottenere che ad espansioni decimali diverse corrispondano numeri reali diversi supponiamo di non rappresentare mai i numeri con periodo nove. Ad esempio scriveremo 0.40000... al posto di 0.3999... Consideriamo la corrispondenza che associa ad ogni coppia $(x,y) \in (0,1)^2$ il numero $f(x,y) \in (0,1)$ ottenuto ponendo $f(x,y) = 0.x_1y_1x_2y_2x_3y_3...$ avendo supposto che $x = 0.x_1x_2x_3...$ e $y = 0.y_1y_2y_3...$. Tale corrispondenza è iniettiva e ciò prova che la potenza di un quadrato è minore di quella del suo lato.²² Per trovare una corrispondenza iniettiva da $(0,1)$ a $(0,1)^2$ possiamo ad esempio considerare la corrispondenza che associa ad ogni $x \in (0,1)$ la coppia $(x, 1/2)$. Per il teorema di Cantor-Bernstein possiamo concludere che $(0,1)$ è equipotente a $(0,1)^2$. $\square$

Questo risultato ai tempi di Cantor appariva stupefacente perché stabiliva l'uguaglianza delle grandezze di due enti geometrici di differente dimensione.

Un modo diverso per convincersi della possibilità di mettere in corrispondenza i punti di una retta ed i punti interni ad un quadrato, è dato dalla considerazione di curve particolari. Una delle curve più famose che progressivamente “riempie” il quadrato è la curva di Peano che si ottiene come limite di una successione di curve definita per ricorsione. La curva di Peano è una funzione continua di $(0,1)$ in $(0, 1)^2$. D'altra parte, stante il fatto

²² Si osservi che la funzione f non è suriettiva poiché, ad esempio, il numero 0.19191919... dovrebbe avere come immagine inversa la coppia 0.11111.... e 0.9999... ma tale coppia non rappresenta un punto interno al quadrato.



che la dimensione di $(0,1)$ è diversa da quella di $(0,1)^2$, non può esistere nessuna funzione bicontinua (continua, invertibile e tale che la sua inversa sia continua) tra $(0,1)$ e $(0,1)^2$.

Problema. Dimostrare che ogni cubo ha la stessa potenza di un suo spigolo.

Dalle proposizioni ora dimostrate segue facilmente il seguente teorema.

Teorema 9.4. L'insieme dei punti del piano ha la potenza del continuo. L'insieme dei punti dello spazio ha la potenza del continuo.

Dim. Poiché R è equipotente a $(0,1)$, $R \times R$ è equipotente a $(0,1) \times (0,1)$. D'altra parte abbiamo già dimostrato che tale insieme è equipotente a $(0,1)$. Quindi $R \times R$ ha la potenza del continuo.

In modo analogo si dimostra che l'insieme dei punti dello spazio ha la potenza del continuo. $\square$

Problema. L'insieme dei triangoli del piano ha potenza maggiore o uguale al continuo?

Problema. L'insieme dei cerchi del piano ha potenza maggiore o uguale al continuo?

Problema. Uno spazio vettoriale di dimensione finita sul campo dei reali quale cardinalità ha?

Problema. Il gruppo delle rotazioni intorno ad un dato punto è ciclico?

Esercizio. Dimostrare che due circonferenze sono equipotenti (Una dimostrazione analitica si ottiene col considerare le equazioni parametriche. Una dimostrazione geometrica si ottiene provando l'equipotenza per circonferenze concentriche e poi uti-

lizzando traslazioni per estendere il risultato ad una qualunque coppia di circonferenze).

Proposizione 9.5. L'insieme dei numeri complessi ha la potenza del continuo.

Dim. Basta ricordare che un numero complesso $x+iy$ è caratterizzato dalla sua parte reale x e dalla sua parte complessa y , cioè da una coppia (x,y) di numeri reali. $\square$

Abbiamo mostrato fino ad ora l'esistenza di tre tipi di insiemi:

- quelli finiti,
- quelli numerabili
- quelli aventi la potenza del continuo.

Ancora una volta sembra difficile trovare qualche cosa di più grande del continuo. Tuttavia il seguente teorema, dovuto a Cantor, mostra che:

dato un insieme S esiste sempre un insieme più grande di S .

Teorema 9.6. (Teorema di Cantor) Dato un insieme S , l'insieme $\mathcal{P}(S)$ delle parti di S ha potenza strettamente maggiore di quella di S . In particolare, l'insieme $\mathcal{P}(\mathbb{R})$ e l'insieme $\mathbb{R}^{\mathbb{R}}$ delle funzioni di variabile reale hanno potenza maggiore del continuo.

Dim. Supponiamo per assurdo che esista una funzione biettiva f di S in $\mathcal{P}(S)$ e consideriamo l'insieme $T = \{x \in S \mid x \notin f(x)\}$. Siccome f è suriettiva esisterà $x_0 \in S$ tale che $f(x_0) = T$. Allora:

- se $x_0 \in T$ avremo che x_0 verifica la proprietà caratteristica di T e quindi $x_0 \notin f(x_0)$ e pertanto $x_0 \notin T$
- se $x_0 \notin T$ allora x_0 non verifica la proprietà caratteristica di T e quindi $x_0 \in f(x_0)$ e $x_0 \in T$.

Si perviene pertanto all'equivalenza $x_0 \in T \Leftrightarrow x_0 \notin T$ che è assurda in quanto afferma che una asserzione coincide con la sua negata.

Per dimostrare la parte rimanente della proposizione osserviamo che $\mathbb{R}^{\mathbb{R}}$ ha ovviamente cardinalità maggiore di $\{0,1\}^{\mathbb{R}}$ e che questo ultimo insieme ha cardinalità uguale a $\mathcal{P}(\mathbb{R})$. $\square$

In definitiva se indichiamo con n un insieme finito con un numero n di elementi, abbiamo visto che esiste una successione

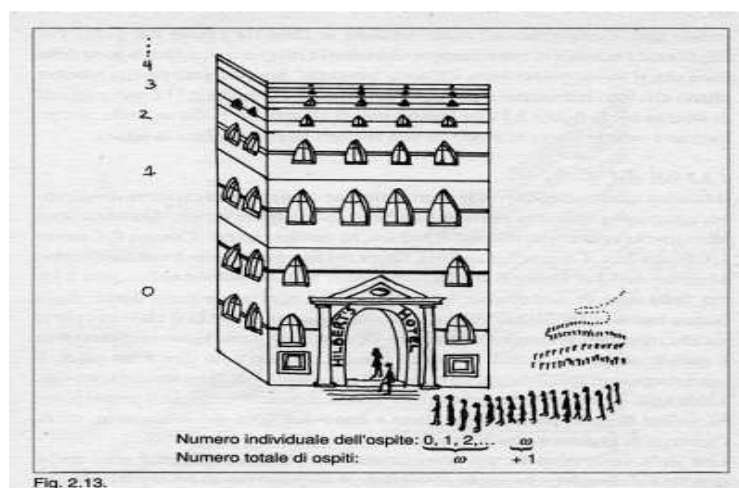
$$1, 2, 3, 4, 5, \dots, n, \mathcal{P}(n), \mathcal{P}(\mathcal{P}(n)), \dots$$

di insiemi che è crescente rispetto alla relazione di “avere potenza minore di”.

LETTURA

Rudy Rucker – L'albergo di Hilbert, da *La mente e l'infinito*, Editore Muzzio, 85-86.

Il famoso matematico David Hilbert, nelle sue conferenze a carattere divulgativo, raccontava spesso la storia di un albergo con infinite stanze: Questo mitico albergo, che chiameremo *Albergo di Hilbert*, ha omega camere: Camera 0, Camera 1, Camera 2, ..., Camera n , e così via. Come nel paragrafo precedente inizieremo a contare da 0. Per fissare le idee, ho disegnato l' Albergo nella figura 2.13. Per farlo stare su una pagina, ho supposto che ogni piano fosse dotato di un fantascientifico condensatore di spazio, uno strumento che fa sì che ogni piano sia alto due terzi del precedente. Anche gli ospiti subiscono la stessa contrazione e quindi, sebbene i soffitti del terzo piano siano alti solo due o tre piedi, il condensatore di spazio rende gli ospiti alti uno o due piedi in modo che si trovino a loro agio. Lascio al lettore come esercizio di dimostrare che, se il primo piano ha soffitti di dieci piedi e ogni piano è alto i due terzi del precedente, allora l'albergo di ω piani è alto 30 piedi.



Una delle caratteristiche più sconcertanti dell'Albergo di Hilbert è che, anche quando è al completo, è sempre possibile trovare posto per nuovi ospiti senza che nessuno debba condividere la sua stanza con un altro!

Supponiamo, per esempio, che ogni camera sia occupata collocando l'ospite n nella Camera n . Supponiamo che a questo punto arrivi un nuovo ospite: l'Ospite a . Dove potrà alloggiare? E' facile! basta chiedere ad ogni cliente di passare nella camera successiva. In tale modo si libera la camera numero 0 e la si può utilizzare per ospitare a . Ma se fosse arrivato un nuovo pulman contenente una quantità numerabile di nuovi turisti. Come collocarli? Facile, basta dire a ciascun ospite della camera n di passare alla camera $2n$. In tale modo tutte le camere di numero dispari rimangono libere. Allora si dice ai turisti del pulmann di collocarsi in tali camere. Precisamente al primo turista che scende dal pulman si assegna la camera $2 \cdot 1 - 1 = 1$, al secondo turista la camera $2 \cdot 2 - 1 = 3$, ... al turista n la camera $2 \cdot n - 1$.

NOTA. Naturalmente le stesse considerazioni le potremmo fare se l'albergo avesse stanze tutte uguali tra loro ma lungo una linea retta infinita (a me piace più immaginarla orizzontale e non verticale forse per questione di vertigini). In questo caso lo spostarsi alla camera successiva sarebbe una vera e propria traslazione. Ciò permetterebbe di interpretare quanto detto in questa lettura dicendo che l'albergo è equiscomponibile con un albergo uguale più una stanza (in generale più n stanze). Il paradosso degli otto raggi esposto nel paragrafo 10 del primo capitolo è analogo salvo il fatto che i punti della successione $A_1, A_2, \dots$ sono stati collocati su di una circonferenza e non su di una retta.

CAPITOLO 5

METODO ASSIOMATICO E STRUTTURALISMO

1. Paradossi e crisi della teoria degli insiemi

Non mi sarei mai iscritto ad una associazione che mi accettasse come mem-bro.

(il comico Groucho Marx).

Maria è destinata ad essere infelice, ama solo le persone che non l'amano.

(da una conversazione tra studentesse)

La scelta di considerare la teoria degli insiemi come base su cui fondare tutta la matematica mostrò ben presto alcune difficoltà di fondo. Dopo pochi anni che la teoria degli insiemi si andava diffondendo cominciarono ad essere scoperti diversi paradossi che mostravano come la teoria degli insiemi fosse contraddittoria. Prima di esporre tali paradossi, esaminiamo un po' più da vicino quelle che sono due caratteristiche principali della nozione di insieme che permettono di distinguerla dalla analoga nozione di classe o da quella di proprietà.

a) Principio di comprensione.

Tale principio afferma che data una qualunque proprietà P , la collezione degli enti verificanti P , che indicheremo con l'espressione $\{x : x \text{ verifica } P\}$, è un insieme. Ad esempio, alla proprietà "essere pari" corrisponde l'insieme dei numeri pari, alla proprietà "essere maggiore di 2" corrisponde l'insieme dei numeri maggiori di due e così via.

b) Principio di sostanzialità.

Tale principio, che è forse quello più importante, afferma che ogni insieme ha carattere di "sostanza" nel senso che, all'interno del discorso scientifico, gli insiemi possono essere trattati allo stesso modo di oggetti individuali come tavolo, atomo, punto.¹ Il principio di sostanzialità comporta che i nomi degli insiemi possono figurare come soggetti all'interno delle frasi e quindi, in particolare, che ha senso dire che un insieme verifica o meno una data proprietà. Ad esempio

"l'insieme dei numeri primi è infinito"

¹ Si veda anche il paragrafo 2 del capitolo precedente.

è una frase in cui è l'intero insieme dei numeri primi a figurare come soggetto e tale insieme verifica la proprietà di essere infinito. Ciò non avviene, ad esempio, nella frase equivalente

"per ogni primo p esiste un primo q maggiore di p ".

D'altra parte per il principio di comprensione ad ogni proprietà è possibile associare un insieme. Pertanto in termini insiemistici possiamo esprimere il principio di sostanzialità dicendo che un insieme può essere elemento di un altro dato insieme. Nell'esempio fatto l'insieme dei numeri primi appartiene all'insieme degli insiemi finiti.

Da osservare anche che il principio di sostanzialità comporta che gli insiemi vengano considerati esistenti indipendentemente dal processo conoscitivo o creativo dell'uomo (come appunto avviene per le "sostanze"). Ciò pone quella che è la più moderna delle teorie matematiche, la teoria degli insiemi, nell'ambito di una delle più antiche filosofie: il platonismo.

Possiamo ora elencare alcuni dei paradossi che nascono dalla teoria degli insiemi. Il primo paradosso si ottiene applicando il principio di comprensione alla proprietà "essere un insieme" e quindi considerando l'insieme universale $U = \{X : X \text{ è un insieme}\}$ di tutti gli insiemi.

Proposizione 1.1. (Paradosso della classe di tutti gli insiemi)

Sia U l'insieme di tutti gli insiemi, allora la cardinalità di U è maggiore della cardinalità di ogni altro insieme. Ciò è in contrasto con il teorema di Cantor che afferma che la cardinalità di U è strettamente minore di quella di $\mathcal{P}(U)$.

Dim. Sia Z un insieme. Allora la funzione $h : Z \rightarrow U$ che associa ad ogni elemento $z \in Z$ il singoletto $h(z) = \{z\}$ è una funzione iniettiva di Z in U . $\square$

Un altro paradosso nasce dal considerare la proprietà $X \notin X$ di essere normale, cioè di non appartenere a se stesso.

Definizione 1.2. Chiamiamo *normale* un insieme che non appartiene a se stesso.

Gli insiemi che usualmente consideriamo sono normali. Ad esempio l'insieme $\{1, \{2, 3\}, \{4, 5, 6\}\}$ non appartiene a se stesso e quindi è normale. L'insieme $\{X : X \text{ è finito}\}$ degli insiemi finiti è

infinito (ad esempio tutti i singoletti $\{n\}$ con $n \in N$ sono finiti) e quindi non appartenendo a se stesso è normale. Invece il suo complemento $\{X : X \text{ è infinito}\}$ essendo infinito appartiene a se stesso e quindi non è normale. Ancora l'insieme di tutti gli insiemi appartiene a se stesso e quindi non è normale. Se tali insiemi sembrano troppo grossi, possiamo considerare l'insieme X degli insiemi che sono descrivibili nella lingua italiana. Tale insieme non è molto grande in quanto, poiché con la lingua italiana posso formulare solo una quantità numerabile di descrizioni, X è numerabile. D'altra parte, avendo definito X tramite la lingua italiana, X appartiene a se stesso.

Problema: Dire se i due insiemi $\{x : x \text{ è un triangolo}\}$ e $\{x : x \text{ non è un triangolo}\}$ sono normali oppure no.

Data la proprietà di essere normale, per il principio di comprensione esisterà un insieme corrispondente, che chiamiamo *insieme di Russell*, cioè l'insieme $Ru = \{X \mid X \text{ normale}\} = \{X \mid X \notin X\}$. D'altra parte, in base al principio di sostanzialità ha senso chiedersi se Ru sia normale o meno. Da queste due osservazioni segue il seguente paradosso dovuto a B. Russell.

Proposizione 1.3. (Paradosso di Russell, 1901) Sia $Ru = \{X \mid X \text{ è normale}\}$ l'insieme di Russell, Allora risulta che

$$Ru \in Ru \Leftrightarrow Ru \notin Ru.$$

Dim. Se $Ru \in Ru$ allora Ru verifica la proprietà definitoria di Ru e quindi $Ru \notin Ru$. D'altra parte se $Ru \notin Ru$ allora Ru verificando la proprietà definitoria di Ru è normale e quindi $Ru \in Ru$. $\square$

Si potrebbe pensare che i paradossi possano nascere solo dalla considerazione di proprietà strane come quella di non appartenere a se stesso o dalla considerazione di insiemi troppo grandi come l'insieme di tutti gli insiemi. In tale caso per evitare guai un matematico deve solo limitarsi a considerare proprietà ed insiemi più familiari e piccoli. Le cose non stanno così, esistono contraddizioni che nascono dalla considerazione di insiemi che appaiono più naturali. Si consideri ad esempio proprietà semplicissime come “avere un numero finito di elementi” oppure “avere un solo elemento”. Allora per il principio di comprensione è possibile considerare gli insiemi $Fin = \{X : X \text{ è finito}\}$ e $Sing = \{X :$

X è un singoletto}. Se si analizza la dimostrazione della Proposizione 1.1 ci si accorge che la stessa dimostrazione permette di provare anche i seguenti paradossi.

Proposizione 1.4. (Paradosso della classe di tutti gli insiemi finiti) Sia la classe Fin di tutti gli insiemi finiti che la classe $Sing$ di tutti i singoletti hanno cardinalità maggiore di ogni altro insieme. In particolare Fin e $Sing$ hanno cardinalità maggiore di $\mathcal{P}(Fin)$ e $\mathcal{P}(Sing)$, rispettivamente. Ciò è in contraddizione con il teorema di Cantor.

La situazione non è molto diversa se si considera la proprietà "avere tre elementi". In base al principio di comprensione esisterà l'insieme corrispondente, cioè l'insieme $T = \{X \mid X \text{ ha tre elementi}\}$, cioè il numero cardinale $[\{a,b,c\}]$. Per tale insieme è possibile provare la stessa strana proprietà di cui gode l'insieme di tutti gli insiemi.

Proposizione 1.5. (Paradosso della classe degli insiemi con tre elementi). L'insieme $T = \{X \mid X \text{ ha tre elementi}\}$ ha più elementi di ogni altro insieme. In particolare T ha più elementi di $\mathcal{P}(T)$ e ciò è in contraddizione con il teorema di Cantor.

Dim. Sia Z un qualunque insieme e consideriamo un qualunque insieme di tre elementi, ad esempio l'insieme $\{a, b, c\}$. Allora la corrispondenza $f : Z \rightarrow T$ che ad ogni elemento $x \in Z$ associa l'insieme $f(x) = \{(a,x), (b,x), (c,x)\}$ è una funzione iniettiva di Z in T . Ciò prova che in T ci sono più elementi che in Z . $\square$

Forse è interessante riformulare il paradosso della classe degli insiemi con tre elementi anche per la classe dei gruppi (classe di cui gli algebristi parlano quotidianamente).

Proposizione 1.6. (Paradosso della classe dei gruppi) L'insieme G dei gruppi ha più elementi di ogni altro insieme. In particolare G ha più elementi di $\mathcal{P}(G)$ (in contrasto con il teorema di Cantor).

Dim. Per prima cosa osserviamo che ogni singoletto $\{z\}$ può essere considerato un gruppo $(\{z\}, +, -, z)$ dove le operazioni sono

definite ponendo $z+z = z$ e $-z = z$. E' infatti un gruppo in cui l'elemento neutro è z e l'inverso di z è ancora z . Sia Z un qualunque insieme. Allora la corrispondenza $f : Z \rightarrow G$ che associa ad ogni $z \in Z$ il gruppo $f(z) = (\{z\}, +, -, z)$ definito dal singoletto $\{z\}$ è iniettiva. Pertanto G ha più elementi di Z . $\square$

In realtà tale proposizione è molto più generale, infatti vale per la classe degli anelli, la classe degli spazi vettoriali e per tutte le classi di strutture definite da un sistema di assiomi.

Proposizione 1.7. Sia T una teoria, allora la classe dei modelli di T conduce ad un paradosso.

Problema: Chiamiamo *club* un insieme qualunque e consideriamo la classe $C = \{x : x \text{ è un club e } io \in x\}$ dei club che mi hanno come membro. Provare che comunque si consideri un insieme Z , C ha cardinalità maggiore o uguale a quella di Z e che quindi C conduce ad un paradosso.

2. Russell, il paradosso del barbiere, Marx e le studentesse

E' importante osservare che nel paradosso di Russell la sola implicazione $Ru \in Ru \Rightarrow Ru \notin Ru$ non risulta paradossale. Per capire tale questione riconsideriamo la battuta di Marx con cui si è iniziato il paragrafo precedente. Detto x l'insieme dei membri di una associazione allora Marx afferma che

$$\forall x (io \in x \Rightarrow io \notin x).$$

Questo non è paradossale ma è solo la prova che $io \notin x$, in altri termini che Marx non può appartenere a nessuna associazione. Detto più in generale, il fatto che una implicazione $A \Rightarrow \neg A$ risulti vera non è affatto paradossale e permette di dedurre solo che A è falsa. Infatti, $A \Rightarrow \neg A$ è logicamente equivalente ad $\neg A$. Per convincersi di questo basta ricordare un po' di logica e che una implicazione del tipo $A \Rightarrow B$ risulta equivalente alla disgiunzione $\neg A \vee B$. Pertanto, in particolare, $A \Rightarrow \neg A$ è equivalente a $\neg A \vee \neg A$ cioè a $\neg A$. D'altra parte le dimostrazioni per assurdo di cui abbiamo parlato nel primo capitolo si basano proprio su tale principio. Risulta invece paradossale una equivalenza del tipo $A \Leftrightarrow \neg A$ che afferma che una asserzione A è sia vera che falsa.

Più vicina alla struttura logica del paradosso di Russell è l'osservazione su Maria che ama solo le persone che non

l'amaro. Dal punto di vista logico viene affermata la validità dell'asserzione $\forall X(ama(Maria, X) \Leftrightarrow \neg ama(X, Maria))$. Ponendo *Maria* al posto di *X* si ottiene in particolare che

$$ama(Maria, Maria) \Leftrightarrow \neg ama(Maria, Maria)$$

che è ovviamente una contraddizione. L'effetto di tale contraddizione è meno catastrofico del paradosso di Russell in quanto prova solo che non può esistere una persona che ama solo le persone che non l'amaro.

Un paradosso simile era noto agli antichi greci sotto la forma del "paradosso del barbiere". In tale paradosso veniva data la seguente definizione di "barbiere"

Definizione 2.1. "Il barbiere è una persona che taglia la barba a tutte e sole le persone che non se la tagliano da sole"

Allora detto Carlo il barbiere davanti al problema di farsi la barba o meno risulta che

- se Carlo non si fa la barba allora, poiché è uno che non si fa la barba da solo, deve farsela (e quindi sbaglia)
- se Carlo si fa la barba allora fa la barba ad uno che se la fa da solo (e quindi sbaglia).

In termini più logici si prova che

$$si_fa_la_barba(Carlo) \Leftrightarrow \neg(si_fa_la_barba(Carlo)).$$

Questa è una contraddizione e prova solo che non può esistere una persona che fa il mestiere indicato dalla Definizione 2.1 cioè che tale definizione non può soddisfatta allo stesso modo come non può essere soddisfatta una definizione del tipo "un barbaranumero è un numero che è sia pari che dispari".

3. Affrontare i paradossi: intuizionismo e metodo assiomatico

Vi furono due modi diversi di reagire alla scoperta di tali paradossi. Quello della maggior parte dei matematici fu semplicemente di fingere di non accorgersene. La teoria degli insiemi risultava essere uno strumento tanto utile che nessuno sembrava disposto a rinunciarci per questioni di rigore scientifico. Quei pochi che invece si occupavano di fondamenti della matematica videro nei paradossi la crisi dell'intero apparato conoscitivo della matematica. Vediamo cosa scrive Frege nel 1903 nei suoi *"Fondamenti dell'aritmetica"*, un testo fondamentale dove per la prima volta si propone una definizione dei numeri interi su base insiemistica.

Nulla di più indesiderabile può capitare ad uno scienziato del fatto che una delle fondamenta del suo edificio si incrina dopo che l'opera è finita. E questa la situazione in cui mi trovo in seguito ad una lettera (contenente il paradosso) inviatami del sig. Bertrand Russell proprio mentre si stava ultimando la stampa di questo volume. . .

"Solatium miseris, socios habuisse malorum". Anch'io ho questo sollievo, se sollievo lo possiamo chiamare; infatti chiunque nelle sue dimostrazioni abbia fatto uso di estensioni di concetti, di classi, di insiemi (compresi i sistemi di Dedekind) si trova nella mia stessa posizione. Non è soltanto questione del mio particolare modo di gettare le fondamenta, ma è in questione la possibilità o meno di dare all'aritmetica un qualsiasi fondamento logico.

I matematici che si occupavano dei fondamenti della matematica tentarono di risolvere la questione dei paradossi in teoria degli insiemi seguendo due vie totalmente diverse tra loro, quella dell'*intuizionismo* e quella dell'*assiomatizzazione* della teoria degli insiemi ed in generale della matematica.

L'intuizionismo. Gli intuizionisti rigettavano completamente la teoria di Cantor e quella parte della matematica che su essa si fondava. Esponenti di questa corrente filosofica, che si ispira a Kant, sono due matematici di rilievo, L.E.J. Brouwer (1881-1966) e H. Weyl (1885-1955). Gli intuizionisti sostengono che la matematica si fonda sull'intuizione *a priori* del tempo e che i suoi enti sono «costruzioni» della mente umana effettuate a partire da questa intuizione. In particolare, l'intuizione del trascorrere del tempo porta all'idea di numero intero che essi ritengono, in comune con i pitagorici, base di tutta la matematica. Non ha senso ricondurre la nozione di numero naturale a quella di insieme o ad un sistema di assiomi perché l'idea di numero naturale è insita in ogni uomo. A partire da tale intuizione, mediante metodi costruttivi di tipo mentale, si doveva poi procedere alla edificazione del rimanente corpo della matematica:

MATEMATICA =
INTUIZIONE DEL NUMERO INTERO +
COSTRUZIONE DEGLI ENTI MATEMATICI

Per gli intuizionisti se una parte della matematica non è ottenibile in tale modo, allora tale parte è da considerare inaffidabile e pertanto deve essere rigettata via.

A partire da tali idee gli intuizionisti sviluppano una matematica ed una logica per molti aspetti diversa e più sottile di quella che usualmente siamo abituati a considerare. Ogni asserzione matematica è giustificata dalla costruzione di una dimostrazione, ogni ente matematico esiste solo se è stato concretamente costruito. Da ciò segue un tipo di logica completamente differente definita grosso modo al modo seguente.

- un'asserzione atomica, cioè priva di connettivi logici, α è valida solo se la costruzione da essa descritta è stata fatta.
- una asserzione $\neg\alpha$ è valida solo se da α è stata derivata una contraddizione
- una asserzione $\alpha\wedge\beta$ è valida solo se si è esibita sia una dimostrazione di α che una dimostrazione di β .
- una asserzione $\alpha\vee\beta$ è valida solo se si è esibita o una dimostrazione di α oppure una dimostrazione di β .
- una asserzione $\exists x(\alpha)$ è valida se e solo se si è costruito un ente matematico verificante la proprietà espressa da α .

Da questa interpretazione seguono leggi della logica notevolmente diverse da quelle classiche. Ad esempio la legge del terzo escluso per gli intuizionisti non è valida perché se valesse $\alpha\vee\neg\alpha$ per ogni asserzione α , allora, qualunque sia l'asserzione saremmo sempre in grado di provare α oppure provare $\neg\alpha$. Ciò potrebbe accadere solo se tutte le questioni della matematica fossero state risolte.

Inoltre la legge di doppia negazione neanche è valida perché il significato dell'asserzione $\neg(\neg\alpha)$ è diverso dal significato α . Ad esempio $\neg\neg\exists x(\alpha)$ significa che è assurdo supporre che non esista un elemento x per cui valga α . Invece $\exists x(\alpha)$ significa che si è in grado di esibire concretamente un elemento che verifica α . Le due cose sono totalmente differenti. Infatti, supponiamo ad esempio che $\exists x(\alpha)$ sia l'affermazione "esiste un numero trascendente". Utilizzando la teoria degli insiemi è possibile provare tale affermazione dimostrando che la cardinalità dell'insieme dei numeri algebrici è numerabile. Pertanto se non esistessero numeri trascendenti l'insieme dei numeri reali sarebbe numerabile e ciò contraddice un teorema di teoria degli insiemi. Stante questa

dimostrazione un intuizionista direbbe che è stata dimostrata l'asserzione $\neg\neg\exists x(\alpha)$ e non l'asserzione $\exists x(\alpha)$. Invece quando, dopo avere definito π , viene dimostrato che π è trascendente (cosa enormemente più complicata), allora è possibile asserire $\exists x(\alpha)$. In altre parole gli intuizionisti utilizzano la negazione in modo che sia possibile distinguere enunciati di teoremi dimostrati per assurdo da enunciati di teoremi dimostrati in modo costruttivo.

L'intuizionismo, per quanto interessante ed affascinante, non ebbe e non ha molti seguaci e normalmente viene ignorato dai matematici contemporanei. Il problema è che l'accettazione dei principi dell'intuizionismo comporta difficoltà enormi nella dimostrazioni dei teoremi. Per mancanza di spazio e per la difficoltà della materia, in questi appunti non ci occuperemo ulteriormente di affascinante modo di vedere la matematica.

Il metodo assiomatico. La presentazione della matematica in forma assiomatica è, come abbiamo già osservato, il merito principale dell'opera di Euclide. Il metodo assiomatico assume però, dalla fine dell'ottocento in poi, aspetti completamente nuovi. L'opera in cui il nuovo modo di considerare il metodo assiomatico si esprime in maniera più completa è *"I fondamenti della geometria"* di David Hilbert la cui prima edizione risale al 1899. Lo scopo che si propone Hilbert è quello di fornire una rigorosa assiomatizzazione degli Elementi di Euclide ma la filosofia che sta alla base di tale libro è totalmente diversa da quella degli Elementi. In Euclide gli assiomi, e quindi i teoremi erano visti come asserzioni vere riguardanti un mondo, quello degli enti matematici, avente una esistenza propria. La matematica si caratterizzava come studio di tali enti (i numeri, i punti, le rette, ...) così come la zoologia si caratterizza per lo studio degli animali e la botanica per lo studio dei vegetali. Totalmente diverso è invece l'atteggiamento che viene assunto dai fautori del moderno metodo assiomatico. La matematica si caratterizza non per avere un particolare oggetto di studio, ma per essere un particolare metodo di indagine quello, appunto, ipotetico-deduttivo. In altre parole:

il metodo assiomatico da strumento di lavoro della matematica viene a coincidere con la matematica stessa.

L'essenza dell'approccio assiomatico consiste nell'esaminare tutte le conseguenze (interessanti) implicite in un particolare

gruppo di ipotesi (gli assiomi), conseguenze ottenute con il solo ausilio di deduzioni senza mai fare riferimento all'esperienza o alla intuizione. È significativo esaminare il diverso atteggiamento rispetto alle definizioni degli enti primitivi. In Euclide gli assiomi sono preceduti dalle definizioni: come si può parlare dei punti e delle rette se non si è prima spiegato che cosa sono i punti e le rette? Le definizioni cercano pertanto, in qualche modo, di indicare ciò di cui si vuole parlare, di facilitare l'intuizione, di distinguere un ente matematico da un altro. Come vedremo nel paragrafo successivo, nell'approccio alla geometria proposto da Hilbert ciò non accade, in un certo senso le definizioni degli enti primitivi mancano del tutto. Semmai si può parlare di nomenclatura; si chiamano, per comodità espositiva, "punti" gli elementi di un dato insieme, "rette" gli elementi di un altro insieme e così via.

4. Un approccio assiomatico alla geometria

I *“Fondamenti della Geometria”* di Hilbert iniziano nel modo seguente.

“Consideriamo tre diversi sistemi di oggetti: chiamiamo "punti" gli oggetti del primo sistema e li indichiamo con $A, B, C, \dots$; chiamiamo "rette" gli oggetti del secondo sistema e li indichiamo con $a, b, c, \dots$; chiamiamo "piani" gli oggetti del terzo sistema e li indichiamo con $\alpha, \beta, \gamma, \dots; \dots$

. . . Noi consideriamo punti, rette e piani in certe relazioni reciproche ed indichiamo queste relazioni con parole come "giacere", "fra", "congruente"; la descrizione esatta e completa, ai fini matematici, di queste relazioni segue dagli assiomi della geometria.”²

² In realtà, già dieci anni prima di Hilbert il matematico e logico Giuseppe Peano aveva enunciato chiaramente il metodo assiomatico in geometria :

“Si ha così una categoria S di enti, chiamati "punti". Questi enti non sono definiti. Inoltre, dati tre punti, si considera una relazione fra essi, indicata con la scrittura $c \hat{I} ab$, la quale relazione non è parimenti definita. Il lettore può intendere col segno S una categoria qualunque di enti, e con $c \hat{I} ab$ una relazione qualunque fra enti di quella categoria;

Ciò comporta tra l'altro che non ha senso definire i punti o le rette, ha piuttosto senso definire una intera classe di strutture geometriche caratterizzata da un particolare sistema di assiomi. Dato poi un modello di un tale sistema di assiomi, saranno chiamati "punti" alcuni suoi elementi e "rette" altri elementi. In altre parole l'attributo di essere punto per un dato oggetto non è definibile in termini di proprietà dell'oggetto stesso (ad esempio essere ciò che non ha parti) ma in termini della sua collocazione all'interno di una struttura verificante un particolare sistema di assiomi. Possiamo anche dire che le definizioni sono date implicitamente dallo stesso sistema di assiomi, ciò nel senso che "punto" è ogni elemento di una struttura verificante una lista di assiomi di carattere geometrico; "numero reale" è ogni elemento di una struttura algebrica verificante un sistema di assiomi per i campi completi archimedei.

Per capire meglio lo spirito del libro di Hilbert e del metodo assiomatico, esponiamo i primi assiomi dei suoi *Fondamenti* che si riferiscono alla geometria piana ed alla relazione di giacenza.³ Se un punto P giace su di una retta r diremo anche che r passa per P .

A1 Per due punti A e B esiste almeno una retta r che passa per A e B .

A2 Per due punti A e B esiste al più una retta che passa per A e B .

A3 Su ogni retta giacciono almeno due punti, ci sono almeno tre punti che non giacciono su di una stessa retta.

Dette poi *parallele* due rette che non hanno nessun punto in comune, consideriamo anche il famoso assioma delle parallele.

avranno sempre valore tutte le definizioni che seguono, e sussisteranno tutte le proposizioni. In dipendenza dal significato attribuito ai segni non definiti S e c $\hat{I} ab$, potranno essere soddisfatti, oppure no, gli assiomi. Se un certo gruppo di assiomi è verificato, saranno pure vere tutte le proposizioni che si deducono.” (in “*Principi di geometria logicamente esposti*”).

³Da notare che una retta non è necessariamente un insieme di punti e che la relazione di "giacenza" non è necessariamente la relazione di appartenenza.

A4 Sia r una retta ed A un punto non giacente in r , allora c'è al massimo una retta passante per A e parallela ad r .

Mostriamo un esempio di teorema geometrico che si può dedurre da tali assiomi.

Teorema 4.1. Tutte le rette hanno lo stesso numero di punti. Un fascio proprio di rette contiene tanti punti quanti sono i punti di una retta più uno.

Dim. Sia r una retta e sia P un punto che non passa per r . Allora ad ogni punto Q di r posso associare la retta per P e per Q . Questa corrispondenza è iniettiva ma l'unica retta che non si può ottenere in tale modo è la parallela ad r per P . In definitiva nel fascio ci sono tanti punti quanti sono quelli di r più 1. $\square$

Nel caso della usuale geometria euclidea tale proposizione mostra che tutte le rette sono equipotenti ed hanno la potenza del continuo e che l'insieme delle rette di un fascio ha la potenza del continuo più 1 (cioè ha la potenza del continuo).

Esistono diversi modelli del sistema A1-A4 di assiomi.

Esempio: il piano della geometria analitica: L'esempio più noto di modello del sistema di assiomi proposto è ottenuto a partire dal campo dei numeri reali. Chiamiamo

- *punto* ogni coppia di numeri reali, cioè ogni elemento di $\mathbb{R}^2$,
- *retta* ogni insieme di punti verificanti una equazione di primo grado del tipo $ax+by=p$ con a, b e p non tutti nulli.

La relazione di *giacenza* coincide (in questo caso) con quella usuale di appartenenza. Da notare che terne proporzionali di coefficienti a, b e c determinano la stessa retta poiché determinano equazioni con le stesse soluzioni. Che per due punti diversi (x_0, y_0) e (x_1, y_1) passa una ed una sola una retta deriva dall'esame del sistema omogeneo

$$ax_0 + by_0 = p$$

$$ax_1 + by_1 = p$$

nelle incognite a, b, p . Infatti bastano alcune nozioni di base di algebra lineare per provare che tale sistema ammette almeno una soluzione non nulla e che tutte le soluzioni sono proporzionali tra loro. Il primo fatto comporta l'esistenza di una retta. Il secondo

fatto che la retta è unica. Non è difficile verificare anche l'assioma delle parallele. Infatti, come è noto, i punti di intersezione di due rette $ax+by = p$ e $a'x+b'y = p'$ si ottengono risolvendo il sistema

$$\begin{cases} ax + by = p \\ a'x + b'y = p' \end{cases}$$

Pertanto le due rette sono parallele se e solo se tale sistema non ammette soluzioni e questo avviene se e solo se i coefficienti a e b sono proporzionali ai coefficienti a' , b' . Allora, dato un punto (x_0, y_0) ed una retta r di equazione $ax+by = p$, trovare una retta passante per (x_0, y_0) e parallela ad r equivale a trovare tre numeri a' , b' , p' tali che

$$\begin{cases} ab' = a'b \\ a'x_0 + b'y_0 = p' \end{cases}$$

Anche in questo caso si tratta di un sistema omogeneo in tre incognite ed anche in questo caso esiste una soluzione e tutte le soluzioni sono proporzionali tra loro. Ciò prova l'esistenza e l'unicità della retta per (x_0, y_0) e parallela ad r .

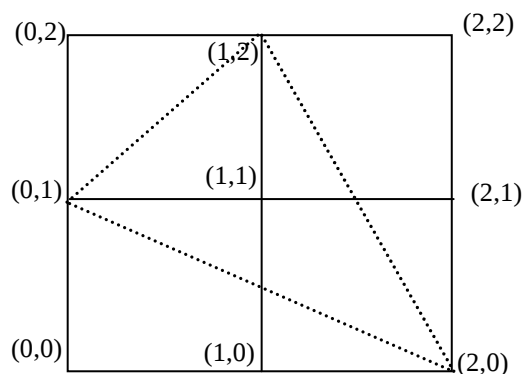
Esempio: Le geometrie finite. Se al campo dei numeri reali sostituiamo un qualunque altro campo, ad esempio il campo degli interi modulo 3, il sistema di assiomi A1-A4 continua ad essere soddisfatto. In tale caso avremo che il modello di geometria è costituito dai nove punti

$$(0,0), (0,1), (0,2), (1,0), (1,1), (1,2), (2,0), (2,1), (2,2),$$

mentre una retta sarà costituita dall'insieme dei punti verificanti una equazione di primo grado. Avremo che le rette passanti per $(0,0)$ saranno le quattro rette di equazione $x = 0$, $y = 0$, e $y = x$ e $y = 2x$, oppure, equivalentemente, $y = -x$. Le rette che non passano per l'origine avranno equazione $ax+by = 1$. Ad esempio una retta è data da

$$\{(x,y) \mid 2x+y = 1\} = \{(0,1), (1,2), (2,0)\}.$$

Procedendo in questo modo non è difficile trovare tutte le rette di questo modello.



La geometria delle carte. L'esempio di piano costruito sul campo degli interi modulo 3 può essere travestito in modo più divertente al modo seguente. Consideriamo nove carte di gioco

che si ottengono prendendo i numeri 1 2 e 3 e i semi di quadri, fiori e picche. Indicheremo tali carte con $1q$, $2q$, $3q$, $1f$, $2f$, $3f$, $1p$, $2p$, $3p$. Diciamo che due carte hanno una proprietà in comune se hanno lo stesso numero oppure lo stesso seme.

Chiamiamo *punto* ognuna di tali carte e chiamiamo *retta* un insieme di tre carte

tale che

- le tre carte hanno lo stesso numero, oppure
- le tre carte hanno lo stesso seme, oppure
- le tre carte non hanno nessuna proprietà in comune.

Allora ad esempio per i punti $1q, 2f$ passa la retta $\{1q, 2f, 3p\}$, per i punti $1q, 1f$, la retta $\{1q, 1f, 1p\}$. È facile rendersi conto che tali rette sono uniche. Inoltre data ad esempio la retta $r = \{1q, 2f, 3p\}$ ed il punto $2q$, la parallela ad r passante per $2q$ è la retta $\{2q, 3f, 1p\}$.

5. Un approccio assiomatico ai numeri reali

In precedenza abbiamo ottenuto i numeri reali partendo dall'insieme dei numeri razionali e “costruendo” la struttura che ci serviva. Hilbert, coerentemente alla sua idea della matematica come sistema ipotetico-deduttivo, diede invece una impostazione assiomatica alla teoria dei numeri reali. In questo paragrafo esponiamo un sistema di assiomi che anche se non coincide con quello proposto da Hilbert è ad esso equivalente e si riferisce alla teoria dei campi ordinati completi. Considereremo in particolare due assiomi: quello di Archimede e quello di completezza (o continuità) che abbiamo già visto nel primo capitolo quando abbiamo definito la nozione di classe di grandezze omogenee. La cosa non deve sorprendere perché la teoria delle grandezze omogenee costituiva per i greci un sostituto della teoria dei numeri reali.

Definizione 5.1. Un elemento x di un campo ordinato si chiama *infinito positivo* se risulta che $x \geq p \cdot 1$ qualunque sia l'intero p . Chiamiamo *infinito negativo* l'opposto di un infinito positivo. Chiamiamo *finito* un elemento che non sia infinito.

L'idea è che x è infinito se, partendo da zero, per quanti passi unitari in avanti si facciano non sia mai possibile raggiungere x . Da notare che x è finito se e solo se esistono p e q tali che $q \cdot 1 \leq x \leq p \cdot 1$, cioè se e solo se x appartiene ad un “intervallo limitato”.

Definizione 5.2. Chiamiamo *infinitesimo positivo (negativo)* un elemento x che sia l'inverso di un elemento infinito positivo (negativo).

Pertanto un elemento δ è un infinitesimo positivo se $\delta > 0$ e risulta che $1/\delta \geq p \cdot 1$ e quindi $\delta \leq 1/(p \cdot 1)$ per ogni naturale p . Possiamo ora definire la nozione di campo archimedeo.

Definizione 5.3. Un campo ordinato $(D, +, \cdot, 0, 1, \leq)$ si dice *archimedeo* se in esso non esistono elementi infiniti positivi.

In altre parole un campo ordinato si dice archimedeo se verifica il seguente assioma:

Assioma di Archimede. $\forall x \exists p \in \mathbb{N} (p \cdot 1 \geq x)$.

Problema. Dimostrare che il campo dei numeri razionali ed il campo dei numeri reali (costruito con le successioni di Cauchy) risultano essere archimedei.

Un campo di razionali non standard è un esempio di campo ordinato non archimedeo. I campi non archimedei sono molto affascinanti perché in essi è possibile sviluppare una teoria degli infiniti e degli infinitesimi. Ad esempio possiamo dire che due elementi x ed y di un campo ordinato sono *infinitamente vicini* se $x - y$ è un infinitesimo. Ovviamente due numeri possono essere infinitamente vicini senza coincidere.

Proposizione 5.4. Se x è un elemento finito allora anche $x + 1$ è finito. Se x è un elemento infinito allora anche $x - 1$ è infinito. Ne segue che l'insieme *Fin* degli elementi finiti non ammette un massimo. Inoltre l'insieme *Inf* degli elementi infiniti di un campo ordinato non ammette minimo.

Dim. Se x è finito allora esiste p tale che $x \leq p \cdot 1$. Ciò comporta che $x + 1 \leq (p + 1) \cdot 1$ e quindi che $x + 1$ è finito. Da ciò segue anche che se x è infinito anche $x - 1$ è infinito. La rimanente parte della proposizione è ovvia. $\square$

Infine arriviamo alla più importante proprietà dei campi ordinati, la completezza. Nel seguito dati due sottoinsiemi A e B di un insieme ordinato scriveremo $A \leq B$ per indicare che ogni ele-

mento di A è minore di ogni elemento di B . In tale caso si dice anche che A e B sono *separati*. Chiameremo “*elemento separatore*” della coppia A e B un elemento u tale $A \leq \{u\} \leq B$ cioè un elemento u maggiore di tutti gli elementi di A e minore di tutti gli elementi di B .

Definizione 5.5. Un campo ordinato si dice *completo* se ogni coppia di insiemi separati ammette elemento separatore.

In altre parole un campo ordinato si dice *completo* se verifica l'assioma:

Assioma di completezza: $\forall A \forall B (A \leq B \Rightarrow \exists u (A \leq \{u\} \leq B))$.

Tutti i campi completi risultano essere archimedei.

Teorema 5.6. Ogni campo completo è archimedeo ma esistono campi archimedei che non sono completi.

Dim. Supponiamo che $(D, +, \cdot, 0, 1, \leq)$ sia un campo completo e consideriamo la coppia Fin ed Inf di sottoinsiemi di D . Se il campo fosse non archimedeo allora Inf sarebbe non vuoto e quindi essendo $Fin \leq Inf$, esisterebbe un elemento separatore u di tale coppia di insiemi. Allora u in quanto maggiorante di Fin è un infinito e quindi appartiene a B . Inoltre u , che è anche mino-
rante di B , appartenendo a B è un minimo di B in contrasto con la Proposizione 5.5.

Il campo dei razionali è un esempio di campo archimedeo non completo. $\square$

Una volta che abbiamo dato le definizioni necessarie possiamo passare alla definizione assiomatica dei numeri reali.

Definizione 5.7. Chiamiamo *campo dei numeri reali* ogni campo completo (e quindi archimedeo).

Ne segue che il campo Q^* dei razionali non standard ed il campo R^* dei reali non standard non essendo archimedei non sono nemmeno completi.

Da notare che, coerentemente con il punto di vista strutturalista, non cerchiamo di definire che cosa sia un numero reale. In-

vece diciamo quando una struttura può essere chiamata “campo dei numeri reali” e chiamiamo “numero reale” ogni elemento di tale struttura. Usualmente si usa l’espressione “il campo dei numeri reali” invece di “un campo dei numeri reali”. Questo modo di dire è giustificato dal seguente teorema, di cui omettiamo la dimostrazione, il quale afferma che, a meno di isomorfismi, esiste un solo campo completo archimedeo.

Teorema 5.8. La teoria dei campi completi è categorica cioè tutti i campi completi sono isomorfi tra loro.

6. Assiomi per eliminare i paradossi dalla teoria degli insiemi

Il metodo assiomatico può sia essere applicato “localmente” a singoli settori della matematica (come la geometria, gli insiemi numerici e così via) oppure può essere utilizzato per tentare di rimettere in piedi la teoria degli insiemi cercando di eliminarne i paradossi. Una volta che si sia trovata una opportuna assiomatizzazione della teoria degli insiemi è possibile poi costruire tutta la matematica come abbiamo fatto nel terzo capitolo.

Naturalmente la questione fondamentale è trovare un sistema di assiomi capace di eliminare i paradossi della teoria degli insiemi che abbiamo trattato nei paragrafi precedenti. Le assiomatizzazioni della teoria degli insiemi che sono state proposte riescono in questo compito limitando opportunamente

- o il principio di comprensione
- oppure quello di sostanzialità.

Limitazione del principio di sostanzialità (la teoria delle classi). La limitazione del principio di sostanzialità si ottiene accettando senza restrizioni il principio di comprensione nel senso che, data una proprietà P , la collezione degli enti verificanti P è sempre una classe. Non sempre le classi hanno però carattere di sostanza e solo le classi per cui ciò avviene hanno il diritto di chiamarsi insiemi. Ora ricordiamo che il carattere di sostanza per una classe si esprime nel fatto che di essa si possa asserire qualche proprietà, e quindi, in termini insiemistici, che essa appartenga ad una classe. Pertanto:

vengono chiamati *insiemi* quelle classi che sono elementi di qualche altra classe.

Detto in breve, x è un insieme se verifica la proprietà $\exists y(x \in y)$. Le classi che non sono insiemi vengono chiamate *classi proprie*.

Delle classi proprie non ha senso dire che verificano o meno una proprietà e non ha senso parlare di cardinalità. Questo modo di procedere fu proposto per la prima volta da von Neumann nel 1925 e poi migliorato da altri. Nella teoria delle classi il principio di comprensione viene riformulato al modo seguente:

$$\exists A \forall X (X \in A \Leftrightarrow X \text{ è un insieme e verifica } P)$$

dove P è una proprietà esprimibile nel linguaggio della teoria degli insiemi ed A è una variabile che non compare in P .

Proposizione 6.1. Nella teoria delle classi il paradosso di Russell non è più valido ma diviene una dimostrazione del fatto che $Ru = \{X : X \text{ è un insieme e } X \notin X\}$ è una classe e non un insieme.

Dim. Nella teoria delle classi l'equivalenza paradossale di Russell $Ru \notin Ru \Leftrightarrow Ru \in Ru$ diventa invece

$$Ru \notin Ru \Leftrightarrow (Ru \text{ non è un insieme oppure } Ru \in Ru).$$

Tale equivalenza equivale a dire che Ru non è un insieme $\square$

Similmente nella teoria delle classi gli altri paradossi citati diventano dimostrazioni del fatto che la classe di tutti gli insiemi, la classe degli insiemi con tre elementi, la classe di tutti i gruppi sono tutte classi proprie.

Una limitazione più forte del principio di comprensione si ottiene sostituendolo con il più debole assioma di isolamento.

Assioma di isolamento: *Dato un insieme S ed una proprietà P , $\{x \in S : x \text{ verifica } P\}$ costituisce un insieme.*

In altri termini l'assioma di isolamento consente solo di "isolare" un opportuno sottoinsieme all'interno di una collezione che sia stata già riconosciuta essere un insieme.

Proposizione 6.2. Se si sostituisce al principio di comprensione l'assioma di isolamento, il paradosso di Russell non è più riproducibile ma si trasforma nella dimostrazione del fatto che, fissato un qualunque insieme S , l'insieme $Ru = \{X \in S \mid X \notin X\}$ non appartiene ad S .

Dim. Se si accetta il principio di isolamento al posto di quello di comprensione allora l'equivalenza paradossale provata da Russell diventa

$$Ru \notin Ru \Leftrightarrow Ru \notin S \text{ oppure } Ru \in Ru.$$

che a sua volta è equivalente a $Ru \notin S$. $\square$

Proposizione 6.3. Se si sostituisce al principio di comprensione l'assioma di isolamento, il paradosso dell'insieme di tutti gli insiemi non è più riproducibile.

Dim. Data la proprietà “essere un insieme” il principio di isolamento permette di definire solo la classe $U_S = \{X \in S : X \text{ è un insieme}\}$ degli insiemi che appartengono ad un dato insieme S . Allora, dato un insieme Z , e $z \in Z$ non è detto che il singoletto $\{z\}$ appartenga a U_S poiché non è detto che $\{z\}$ appartenga ad S .

Proposizione 6.4. Se si sostituisce al principio di comprensione l'assioma di isolamento, il paradosso dell'insieme degli insiemi con tre elementi non è più riproducibile.

Dim. Anche in questo caso il principio di isolamento permette di asserire solo che, dato un insieme S , la proprietà “avere tre elementi” definisce un insieme $T = \{X \in S \mid X \text{ ha tre elementi}\}$. Dato allora un insieme Z qualsiasi ed $z \in Z$, non è detto che $\{(a,z), (b,z), (c,z)\}$ appartenga ad S e quindi a T . $\square$

Proposizione 6.5. Tutti i paradossi della teoria degli insiemi che abbiamo esposto si traducono in una dimostrazione che la classe U di tutti gli insiemi non è un insieme.

Dim. Sia U la classe di tutti gli insiemi, se U fosse un insieme allora l'assioma di isolamento coinciderebbe con l'assioma di comprensione (basterebbe riferirsi sempre ad S come all'insieme in cui isolare l'insieme voluto). Ma dall'assioma di comprensione genera paradossi, e ciò prova che U non è un insieme. $\square$

7. La teoria di Zermelo-Fraenkel

La limitazione al principio di comprensione è presente nella più diffusa teoria degli insiemi che prende il nome di teoria Zermelo-Fraenkel. In tale teoria assumiamo come unici concetti primitivi

la nozione di *insieme* e la *relazione di appartenenza* che indichiamo. Pertanto un modello della teoria di Zermelo-Fraenkel sarà costituito da un insieme con una relazione binaria $\in$.

Nel seguito elenchiamo alcuni degli assiomi di tale teoria senza nessuna pretesa di rigore.

A1. Assioma dell'estensionalità: due insiemi che hanno gli stessi elementi sono uguali. In breve

$$[\forall z(z \in x \Leftrightarrow z \in y)] \Rightarrow x = y$$

Naturalmente è evidente che $x = y \Rightarrow \forall z(z \in x \Leftrightarrow z \in y)$, pertanto l'assioma di estensionalità può essere riformulato al modo seguente:

$$x = y \Leftrightarrow [\forall z(z \in x \Leftrightarrow z \in y)].$$

Tale principio afferma che due insiemi che hanno gli stessi elementi sono uguali, in altre parole due insiemi che hanno la stessa "estensione" coincidono. Ciò differenzia un insieme dal procedimento usato per descriverlo (cioè dalle proprietà che lo definiscono). Per fare un esempio, anche se le proprietà "essere pari" e "terminare per 0, 2, 4, 6 o 8" sono diverse (pur essendo equivalenti) i due insiemi $\{x \in N \mid x \text{ è pari}\}$ e $\{x \in N \mid \text{l'ultima cifra di } x \text{ è } 0, 2, 4, 6 \text{ o } 8\}$ sono uguali. La portata del fatto che in teoria degli insiemi si assume il punto di vista estensionale è più evidente se si considera il concetto di funzione. Come è noto in teoria degli insiemi si accetta la seguente definizione di funzione:

- chiamiamo *funzione* di N in N una relazione binaria $\mathcal{R}$ in N tale che per ogni $x \in N$ esiste uno ed un solo y tale che $(x, y) \in \mathcal{R}$.

Un modo diverso di procedere, che a volte si trova nei libri di scuola, potrebbe essere il seguente:

- chiamiamo *funzione* una *legge* o un *processo* che permette di ottenere un numero y una volta che si sia fissato un numero x .

Una tale definizione appare più intuitiva ma meno rigorosa. Infatti è poco chiaro che cosa si debba intendere per "legge" o "processo". Inoltre in generale noi accettiamo che, ad esempio, equazioni come

$$y = (x-1)^2 + 2x \quad \text{e} \quad y = x^2 + 1$$

rappresentano la stessa funzione pur rappresentando processi di calcolo diversi. Il fatto che le due funzioni siano la stessa significa che quello che interessa è il fatto che se si fissa x in entrambi i

casi si ottiene lo stesso y . In altre parole quello che interessa è l'estensione $\{(x,y) \in R \times R : y = (x-1)^2 - x^2\}$ della proprietà $y = (x-1)^2 - x^2$ che, di fatto, coincide con l'estensione $\{(x,y) \in R \times R : y = x^2 + 1\}$ della proprietà $y = x^2 + 1$. Pertanto l'idea di funzione è "estensionale"; quello che interessa è solo la "tabella" (l'insieme delle coppie) non il modo di calcolare tale tabella.

Un tale punto di vista, tipico dei matematici, potrebbe non essere accettato da un informatico che tiene anche conto dei tempi di calcolo di una funzione. In tale caso è evidente che il punto di vista estensionale non è più valido e che $y = (x-1)^2 - x^2$ e $y = x^2 + 1$ denotano cose diverse (cioè algoritmi diversi).

A2. Assioma dell'insieme vuoto: esiste un insieme che non ha elementi e che chiameremo *insieme vuoto*:

$$\exists x(\forall z(\neg z \in x)).$$

Dall'assioma di estensionalità segue che esiste un solo insieme vuoto che indicheremo con $\emptyset$.⁴

A3. Assioma del singoletto. Dato un insieme x esiste un insieme z che ha x come unico elemento:

$$\forall x \exists y \forall z (z \in y \Leftrightarrow x = z)$$

Per l'assioma di estensionalità l'insieme z è unico e lo denoteremo con $\{x\}$. Questi primi assiomi permettono di asserire l'esistenza di un buon numero di insiemi. Infatti siamo sicuri dell'esistenza degli insiemi $\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \dots$

⁴ Da A1 e A2 segue che la teoria parla solo di insiemi e non di "elementi". Infatti se accettassimo la presenza di un "elemento" a che non contiene a sua volta elementi, allora a dovrebbe coincidere con l'insieme vuoto. Ciò significa che non ha senso considerare come insieme un insieme di oggetti che non siano a loro volta insiemi. Ad esempio l'insieme delle monete che ho in tasca non rientra in tale teoria. Questo rende il tutto poco naturale ma evita molte complicazioni. Un modo diverso di procedere sarebbe accettare che esistano due tipi di oggetti matematici: gli insiemi e gli elementi che non sono insiemi. Per gli elementi che non siano insiemi non vengono imposti i due assiomi A1 e A2.

Un modello dei primi tre assiomi. Nello spirito del metodo assiomatico non dobbiamo necessariamente interpretare gli oggetti che abbiamo chiamato “insiemi” e la relazione che abbiamo chiamato di “appartenenza” nel modo usuale. In effetti per avere una interpretazione si deve avere una collezione di elementi più una relazione binaria. Ad esempio se chiamiamo *insieme* un numero naturale e poniamo $n \in m$ se m è il successivo di n allora abbiamo che gli assiomi A1, A2 ed A3 sono soddisfatti. Infatti A1 segue dal fatto che la funzione successore è iniettiva. A2 è vera perché 0 non è successore di nessun elemento. A3 è verificata perché ogni numero naturale ha uno ed un solo successore. Non è invece verificato il seguente assioma.

A4. Assioma dell'unione di due insiemi. Dati due insiemi x ed y esiste un insieme z che ha come elementi o gli elementi di x oppure gli elementi di y .

$$\forall x \forall y \exists z (\forall z' (z' \in z \Leftrightarrow z' \in x \vee z' \in y))$$

L'insieme z per l'assioma di estensionalità è unico e viene indicato con $x \cup y$. Da tale assioma e da A3 segue che dati due insiemi x_1 e x_2 esiste l'insieme $\{x_1, x_2\} = \{x_1\} \cup \{x_2\}$ caratterizzato dal contenere solo gli elementi x_1, x_2 . Più in generale, dati gli insiemi $x_1, \dots, x_n$ esiste l'insieme $\{x_1, \dots, x_n\} = \{x_1\} \cup \dots \cup \{x_n\}$ che ha $x_1, \dots, x_n$ come unici elementi. Da notare che i numeri naturali non verificano A4 in quanto dati x ed y , con $x \neq y$, non è vero che esiste un z che è successore sia di x che di y .

L'assioma dell'unione di due insiemi può essere esteso in modo da ottenere l'unione di una classe qualunque di insiemi.

A5. Assioma dell'unione degli elementi di una classe. Di ogni classe di insiemi è possibile fare l'unione:

$$\forall x \exists y \forall z (z \in y \Leftrightarrow \exists z' (z' \in x \wedge z \in z')).$$

Per l'assioma di estensionalità l'insieme y di cui si dichiara l'esistenza è unico. Nel seguito lo denoteremo con $\cup x$ oppure $\cup \{x' : x' \in x\}$ e parleremo di *unione degli elementi di x* .

Dati due insiemi x ed y diremo che x è una *parte di* y e scriveremo $x \subseteq y$ se ogni elemento di x è anche un elemento di y , cioè se vale l'implicazione

$$z \in x \Rightarrow z \in y.$$

A6. Assioma dell'insieme delle parti: per ogni insieme x esiste un insieme y costituito da tutte e sole le parti di x :

$$\forall x \exists y \forall z (z \in y \Leftrightarrow z \subseteq x).$$

Per l'assioma di estensionalità l'insieme y è unico. Denoteremo tale insieme con $\mathcal{P}(x)$.

Il prossimo assioma postula l'esistenza di un insieme infinito. Per fare questo dobbiamo definire la nozione di funzione e poi quella di equipotenza.

Definizione 7.1. Dati due insiemi x ed y chiamiamo *coppia* di elementi x ed y l'insieme $\{x, \{x, y\}\}$.

Che tali insieme esista segue dal fatto che è l'unione degli insiemi $\{x\}$ ed $\{\{x, y\}\}$. Segue la solita definizione di prodotto cartesiano, di relazione e di funzione. Ciò permette di definire la nozione di equipotenza tra due insiemi.

Definizione 7.2. Chiameremo *infinito* un insieme che sia equipotente ad una sua parte propria.

A7. Assioma dell'infinito: esiste un insieme infinito.

A questo punto possiamo introdurre l'assioma di isolamento. Tale assioma coinvolge la nozione di *proprietà*. Dobbiamo allora dare una definizione rigorosa di tale concetto. Tale definizione è possibile solo nell'ambito della logica matematica e quindi per poterla dare ci vorrebbe troppo tempo. Qui ci limitiamo a dire che intendiamo riferirci a tutte le proprietà che possono essere descritte nel linguaggio della teoria degli insiemi, cioè utilizzando solo il simbolo di appartenenza $\in$ ed eventualmente i connettivi logici “non”, “esiste”, “per ogni”, “oppure”, “e”.

A8. Assioma di isolamento: sia p una proprietà nel linguaggio della teoria degli insiemi, allora per ogni insieme z esiste un insieme z' tale che

$$x \in z' \Leftrightarrow x \in z \text{ ed } x \text{ verifica } p.$$

L'insieme z' viene indicato anche con $\{x \in z : x \text{ verifica } p\}$.

Per completezza enunciamo anche l'ultimo assioma della teoria ZF che non è molto intuitivo ma molto utile per motivi tecnici. Nel seguito chiamiamo *funzionale* una proprietà $p(x,y)$ scritta nel linguaggio della teoria degli insiemi tale che per ogni x esiste un unico y per cui $p(x,y)$ sia vera. In altre parole una proprietà funzionale è una proprietà a due posti che quando venga interpretata definisce una funzione. Data una proprietà funzionale $p(x,y)$ ed un insieme z chiamiamo *immagine di z tramite $p(x,y)$* un insieme z' tale che per ogni $x \in z$ esiste $x \in z'$ tale che $p(x,y)$ sia vera. Il prossimo assioma assicura che l'immagine di un insieme tramite una proprietà funzionale è ancora un insieme.

A9. Assioma di sostituzione: Sia $p(x,y)$ una proprietà funzionale. Allora per ogni insieme z esiste un insieme z' che è immagine di z tramite $p(x,y)$.

8. Assioma della scelta

L'elenco dei possibili assiomi per la teoria degli insiemi non si esaurisce con quelli che abbiamo dato. Ad esempio è di particolare importanza l'assioma della scelta senza il quale molti teoremi della matematica non possono essere provati. Per poterlo enunciare diamo la seguente definizione:

Definizione 8.1. Se Σ è un insieme non vuoto costituito da insiemi non vuoti, si chiama *funzione di scelta* su Σ ogni applicazione $f: \Sigma \rightarrow \bigcup_{X \in \Sigma} X$ tale che $f(X) \in X$ per ogni $X \in \Sigma$.

La funzione di scelta in un certo senso “sceglie” in ogni insieme $X \in \Sigma$ un suo elemento $f(X)$. Ad esempio se Σ è un insieme di sottoinsiemi di N , allora una funzione di scelta potrebbe consistere nell'associare ad ogni insieme X in Σ il minimo elemento di X .

A10. (Assioma della scelta). Ogni classe Σ di insiemi non vuoti ammette una funzione di scelta.

Un modo ovviamente equivalente di formulare tale assioma è di utilizzare la nozione di famiglia invece che quella di classe. Infatti, data una famiglia $(A_i)_{i \in I}$ di insiemi non vuoti possiamo chiamare *funzione di scelta* ogni funzione $f: I \rightarrow \bigcup_{i \in I} A_i$ tale che $f(i) \in A_i$.

A10'. Per ogni famiglia $(A_i)_{i \in I}$ di insiemi non vuoti esiste una funzione di scelta.

La teoria che si ottiene aggiungendo a ZF tale assioma viene denotata con ZFC . E' difficile capire perché si debba assumere come assioma una cosa tanto evidente ed intuitiva. Il fatto è che la nostra intuizione si basa su classi finite di insiemi ed in tale caso l'assioma della scelta è conseguenza degli altri assiomi.

Proposizione 8.2. Se si considera una classe Σ che sia finita, allora l'assioma della scelta può essere dedotto dagli assiomi di ZF .

Dim. Infatti se $\Sigma = \{X_1, X_2, \dots, X_n\}$ è un insieme finito di insiemi non vuoti, allora esistono $x_1 \in X_1, x_2 \in X_2, \dots, x_n \in X_n$. Pertanto esistono le coppie $(X_1, x_1), (X_2, x_2), \dots, (X_n, x_n)$. Per l'assioma della coppia possiamo asserire che esiste anche l'insieme $\{(X_1, x_1), (X_2, x_2), \dots, (X_n, x_n)\}$ che è la funzione di scelta cercata. $\square$

Proviamo a fare un passo in avanti e supponiamo di avere una classe infinita Σ di insiemi non vuoti. In tale caso, è ancora possibile "scegliere" un elemento x in ciascun X in Σ . Tuttavia questa volta nulla ci assicura che esista un insieme costituito da tutte le coppie (X, x) che si sono ottenute con queste infinite scelte. In altre parole niente ci assicura che la totalità di tali coppie costituisca una funzione. Infatti quando fu proposto per la prima volta esso ha suscitato molte perplessità e reazioni negative. Infatti si aveva la tendenza a credere che un insieme esiste solo se si riescono a descrivere i suoi elementi tramite una proprietà comune. Similmente si era portati a pensare che una funzione esiste solo se è possibile definire in modo preciso quale è la legge che associa l'input all'output. Allora chi ci assicura che esiste sempre un criterio generale che permetta di scegliere all'interno di ogni insieme in Σ un suo elemento?

Scarpe e calzini. Per ben chiarire questo fatto è utile e divertente illustrare l'esempio di Bertrand Russel. Supponiamo preliminarmente di avere una famiglia infinita $(S_i)_{i \in I}$ di paia di scarpe (possiamo visualizzare S_i come una scatola su cui è attaccata una etichetta i per distinguerla dalle altre) e di voler individuare una

“funzione di scelta” che permetta di scegliere in ogni S_i una scarpa $f(i) \in S_i$ particolare. In questo caso non ci sono difficoltà, possiamo scegliere in ogni paio, ad esempio, la scarpa destra. Supponiamo invece di disporre una famiglia infinita $(C_i)_{i \in I}$ di paia di calzini. Allora poiché i due calzini di un dato paio non sono distinguibili, non sarà semplice definire una funzione di scelta che consenta di scegliere in ogni paio C_i un particolare calzino. Il procedere “scegliendo a caso” non sembra fare parte dell’universo matematico. L’assioma della scelta afferma che una tale funzione di scelta esiste e questo indipendentemente dal fatto che si sia capaci di descriverla.

Il seguente teorema mostra come, a volte, facciamo uso di questo postulato senza rendercene conto.

Teorema 8.3. L’assioma della scelta è equivalente a dire che il prodotto cartesiano di insiemi non vuoti è non vuoto. In altri termini è equivalente a dire che, data una qualunque famiglia di insiemi $(A_i)_{i \in I}$:

$$A_i \neq \emptyset \quad \forall i \in I \Rightarrow \times_{i \in I} A_i \neq \emptyset.$$

Dim. Basta osservare che la definizione di prodotto cartesiano di una famiglia di insiemi (si veda l’appendice) è proprio che $\times_{i \in I} A_i$ coincide con l’insieme di tutte le funzioni di scelta per la famiglia $(A_i)_{i \in I}$.

Naturalmente, stante la Proposizione 8.2, per poter dimostrare che il prodotto cartesiano di un numero finito di insiemi non vuoti è non vuoto non è necessario l’assioma della scelta.

Va sottolineato che molti altri teoremi, di fondamentale importanza nella teoria degli insiemi, sono equivalenti all’assioma della scelta. Tra i più famosi ricordiamo:

- **Lemma di Zorn.** *Se $(S, \leq)$ è un insieme parzialmente ordinato tale che ogni catena ha maggiorante, allora $(S, \leq)$ ha un elemento massimale.*
- **Principio del buon ordinamento.** *In ogni insieme X può essere definito un buon ordinamento.*

9. Dimostrare o confutare l'assioma della scelta

Si pone il problema se l'assioma della scelta sia in realtà un teorema della teoria degli insiemi. Ad esempio, König riteneva di avere dimostrato che l'assioma della scelta non può essere vero.

Teorema 9.1. (König) L'assioma della scelta è falso.

Dim. Supponiamo per assurdo che tale assioma sia vero e che quindi risulti che ogni insieme ammette un buon ordinamento. Allora in particolare il campo R dei numeri reali ammette un buon ordinamento (che naturalmente non coincide con l'ordinamento usuale tra numeri reali). Riprendendo il discorso fatto quando abbiamo esposto il paradosso di Berry, chiamiamo “definibile” un numero reale r se battendo sulla tastiera del mio computer un certo numero di caratteri ottengo una descrizione in lingua italiana di tale numero. Ad esempio la radice positiva di 2 è definibile poiché è l'unico numero reale che soddisfa la condizione $(x^2=2) \wedge (x>0)$. Consideriamo ora l'insieme dei numeri reali che sono definibili. Per quanto osservato nel capitolo 6 nel paragrafo sulla definibilità e la numerabilità, tale insieme è numerabile. Pertanto, non essendo R numerabile, l'insieme dei numeri non definibili ha la potenza del continuo e quindi non è vuoto. Consideriamo ora la frase “il più piccolo numero non definibile” in cui ci si riferisce al buon ordinamento che abbiamo supposto esistente. Tale frase è ovviamente una definizione in italiano di un numero $\underline{r}$ che risulta quindi definibile. D'altra parte, poiché il minimo di un insieme appartiene all'insieme, $\underline{r}$ non è descrivibile e ciò è un assurdo. $\square$

Non è chiaro dove sia l'errore di tale dimostrazione tuttavia è certo che essa non è accettabile. Infatti lo schema adottato è lo stesso del paradosso di Berry e quindi se accettassimo la sua validità dovremmo accettare anche che l'ordinamento usuale in N non sia di buon ordinamento. Questa cosa sarebbe difficile da digerire perché equivarrebbe ad ammettere l'esistenza di una successione infinita di numeri interi positivi sempre più piccoli.

In realtà l'assioma della scelta è indipendente dagli altri assiomi della teoria degli insiemi, cioè non può essere né provato né confutato. Infatti vale il seguente teorema la cui dimostrazione, che richiede nozioni avanzate di logica matematica, omettiamo.

Teorema 9.2. L'assioma della scelta è indipendente dai rimanenti assiomi.

10. Ipotesi del continuo

Un altro assioma particolarmente importante è l'ipotesi del continuo che abbiamo già incontrato nel capitolo precedente. Abbiamo visto che la potenza di $\mathcal{P}(N)$ è strettamente maggiore di quella di N . Appare allora naturale chiedersi se esistono insiemi che hanno potenza maggiore del numerabile e minore di quella di $\mathcal{P}(N)$. Cantor formulò l'ipotesi per cui non esiste un tale insieme. Trasformiamo questa ipotesi in un assioma.

A11. Ipotesi del continuo. Non esiste un insieme che abbia potenza strettamente maggiore del numerabile e strettamente minore del continuo.

Detto in altri termini l'ipotesi del continuo afferma che nel piano euclideo esistono solo insiemi finiti, numerabili o aventi la potenza del continuo. Per molti anni nessun matematico ha saputo confutare o provare l'ipotesi del continuo. Nel 1963 il logico matematico Cohen ha dato una risposta a tale ipotesi mostrando che è indipendente dai rimanenti assiomi della teoria degli insiemi.

Teorema 10.1. (Teorema di Cohen). L'ipotesi del continuo è indipendente dai rimanenti assiomi.

In altre parole l'ipotesi del continuo non può né essere provata né essere confutata a partire dai sistemi di assiomi che usualmente vengono accettati per la teoria degli insiemi. Più in generale è stata provata anche l'indipendenza della seguente asserzione.

Ipotesi generalizzata del continuo. Dato un insieme infinito S non esiste un insieme che abbia potenza strettamente maggiore di S e strettamente minore di $\mathcal{P}(S)$.

L'ipotesi generalizzata del continuo può essere espressa anche utilizzando la teoria dei numeri cardinali.

Ipotesi generalizzata del continuo. Per ogni cardinale x non esiste nessun cardinale tra x e 2^x .

Se si accetta tale ipotesi i cardinali possono essere ordinati nella seguente successione crescente:

$$1 < 2 < 3 < \dots < \aleph_0 < 2^{\aleph_0} < 2^{2^{\aleph_0}} \dots$$

ed tra due elementi della successione non è possibile trovare nessun altro cardinale.

In definitiva:

risultano indipendenti dal sistema-base degli assiomi per la teoria degli insiemi sia l'ipotesi del continuo che l'assioma della scelta.

Si ripresenta allora per la teoria degli insiemi la stessa situazione dell'assioma delle parallele in geometria. Così come coesistono geometrie euclidee e geometrie non euclidee, è possibile sviluppare sia teorie degli insiemi in cui vale l'ipotesi del continuo che teorie in cui vale l'ipotesi opposta.

Da notare che, oltre che dai paradossi, che si spera possano essere eliminati da una opportuna assiomatizzazione, il ruolo della teoria degli insiemi viene messo in crisi ancora di più dalle dimostrazioni di indipendenza. Abbiamo già visto che l'indipendenza dell'assioma delle parallele in geometria con il mostrare la possibilità di più geometrie aveva tolto il ruolo centrale che la geometria euclidea aveva sempre assunto. Allo stesso modo, l'indipendenza dell'ipotesi del continuo con il mostrare la possibilità di più teorie degli insiemi diverse tra loro poneva il problema di quale fosse quella giusta e toglieva quindi il carattere assoluto che si pensava avesse l'intuizione della nozione di insieme.

Problema. Sia X un insieme numerabile di punti di un foglio di carta e supponiamo che una macchia di inchiostro cada sul foglio in modo da coprire X . Dire quale è la cardinalità della macchia di inchiostro.

11. Categoricalità, completezza, consistenza, indipendenza

In questo paragrafo vogliamo elencare alcune possibili proprietà per teorie. Le definizioni che daremo non saranno rigorose poiché solo nell'ambito di un approccio formale alla logica matematica lo potrebbero essere. Nel seguito indicheremo con T un si-

stema di assiomi per una teoria, indicheremo con α una asserzione e con l'espressione $T \vdash \alpha$ il fatto che α sia un teorema di T . Elenchiamo alcune caratteristiche che può avere una teoria. La prima è la categoricità che abbiamo già visto per le terne di Peano e per la teoria dei campi completi.

Definizione 11.1. Una teoria T viene detta *categorica* se tutti i suoi modelli sono isomorfi tra loro.

Quando una teoria è categorica i matematici a volte dicono anche che “esiste un unico modello a meno di isomorfismi”.

Esempi: La teoria dei gruppi non è categorica (ammette per esempio sia modelli finiti che modelli infiniti). La teoria degli anelli non è categorica (basti pensare all'anello $\mathbb{Z}$ degli interi ed all'anello $\mathbb{Q}$ dei numeri razionali. La teoria degli insiemi ordinati non è categorica (l'intervallo $[0,1]$ e l'intervallo aperto $(0,1)$ non sono isomorfi perché il primo contiene un minimo ed il secondo no).

Esempi: La teoria dei gruppi con cinque elementi è categorica. Il sistema di assiomi della geometria euclidea è categorico.

Definizione 11.2. Una teoria T viene detta *consistente* se non esiste nessuna asserzione α tale che $T \vdash \alpha$ e $T \vdash \neg\alpha$. T viene detta *soddisfacibile* se ammette un modello, cioè una struttura matematica che verifica tutti gli assiomi in T .

Il modo più semplice per provare che una teoria è consistente è mostrare che è soddisfacibile. Infatti se la teoria ammette un modello M e se esistesse per assurdo α tale che $T \vdash \alpha$ e $T \vdash \neg\alpha$, allora in tale modello l'asserzione α dovrebbe essere sia vera che falsa.

Esempi: La teoria dei gruppi è consistente. Infatti se la teoria dei gruppi fosse inconsistente allora, preso ad esempio il gruppo additivo degli interi relativi $(\mathbb{Z}, +)$, in tale gruppo dovrebbe valere sia α che $\neg\alpha$. Un esempio di teoria inconsistente è dato dalla teoria degli insiemi, come mostrano le antinomie. D'altra parte tutte le teorie studiate in algebra sono consistenti (altrimenti non avrebbe senso studiarle).

Esempi: Un esempio di teoria inconsistente è la teoria T degli insiemi ordinati finiti che non hanno elementi minimali. Infatti se un insieme ordinato $(S, \leq)$ non ha elementi minimali, allora preso un elemento $x \in S$ esiste un $x_1 < x$. Poiché x_1 non può essere minimale allora esiste $x_2 < x_1$, ... procedendo in questo modo si prova l'esistenza di una successione infinita di elementi x_n di S tutti diversi tra loro. Ciò prova che esistono infiniti elementi in S . Pertanto in T può essere dimostrata sia l'asserzione "esistono infiniti elementi" che l'asserzione "non esistono infiniti elementi".

Definizione 11.3. Una teoria viene detta *indipendente* se non capita mai che un assioma $\alpha \in T$ possa essere provato dai rimanenti assiomi $T - \{\alpha\}$.

In generale le teorie che si studiano sono, oltre che consistenti, anche indipendenti. Infatti se un assioma fosse dimostrabile a partire da altri, allora lo si potrebbe semplicemente cancellare senza alterare la teoria stessa.

Definizione 11.4. Data una teoria T diciamo che una asserzione α è *indipendente* da T se non accade né che $T \vdash \alpha$ né che $T \vdash \neg\alpha$ cioè α non può essere né provata né confutata a partire dagli assiomi di T .

Definizione 11.5. Una teoria T è detta *completa* se ogni asserzione si può o provare o confutare in T .

Equivalente mente possiamo chiamare completa una teoria tale che non esistono asserzioni che non dipendano da T . Da notare che tale nozione di completezza, che si riferisce ad un sistema di assiomi, è completamente diversa dalla nozione di completezza per un campo ordinato che si riferisce alla relazione d'ordine in una particolare struttura. Da notare ancora che in logica matematica si dimostra che ogni teoria categorica T è anche completa ma che il viceversa non vale.

Esempio: Sia T l'insieme degli assiomi della geometria euclidea tranne l'assioma delle parallele e sia α l'assioma delle parallele. Allora abbiamo visto che α è indipendente da T .

Esempio: Ad esempio sia T la teoria dei gruppi ed α la proprietà commutativa, allora α non può essere dimostrata (perché in tale caso tutti i gruppi sarebbero commutativi). D'altra parte nemmeno $\neg\alpha$ può essere dimostrata (perché in tale caso tutti i gruppi sarebbero commutativi). Pertanto α è indipendente da T . Questo mostra anche che la teoria dei gruppi non è completa.

12. Tre diverse ideologie per il metodo assiomatico

Fino ad ora abbiamo esposto esempi di applicazione del metodo assiomatico che sono stati proposti allo scopo di trovare un fondamento sicuro alla geometria, all'analisi (tramite la teoria dei numeri reali) ed alla teoria degli insiemi. Chiameremo *fondazionale* un tale modo di utilizzare il metodo assiomatico. È necessario però specificare che tale metodo attualmente viene utilizzato con uno spirito completamente diverso ed in chiave che viene detta *strutturalista* per il ruolo importante che gioca la nozione di struttura matematica.⁵ Nel seguito esporremo man mano le idee dello strutturalismo evidenziandone le differenze con l'approccio fondazionale di Hilbert e con il metodo presente negli *Elementi* di Euclide. Ci riferiamo a tale scopo ad alcuni aspetti fondamentali della matematica.

Creazione o sistemazione

Hilbert e Euclide. Il punto di vista fondazionale di Hilbert e l'approccio assiomatico di Euclide non hanno lo scopo di “creare” nuova matematica, di trovare nuovi risultati. Piuttosto si concentrano solo sulla organizzazione logica (Euclide) e sulla giustificazione (fondazione) della matematica già esistente. Infatti i *Fondamenti della Geometria* di Hilbert rappresentano un rifacimento, senza imperfezioni dal punto di vista assiomatico, della geometria di Euclide, rifacimento che si mantiene abbastanza vicino all'originale. Gli *Elementi* di Euclide sono essenzialmente una esposizione logicamente ordinata dell'insieme dei teoremi conosciuti all'epoca.

Strutturalisti. Per la scuola strutturalista il metodo assiomatico viene invece utilizzato per cogliere la struttura comune a diversi

⁵ Tra i fautori più radicali dello strutturalismo, ricordiamo il gruppo di matematici francesi che si autodefiniscono Bourbakisti e che hanno scritto una notevole serie di libri tutti firmati con il nome di un matematico inesistente che hanno chiamato Bourbaki.

rami della matematica allo scopo di unificare trattazioni diverse. È una scelta dettata da questioni che potremmo dire di organizzazione del lavoro scientifico e non da esigenze di fondazione o di analisi critica. Ad esempio gli assiomi della teoria dei gruppi collegano proprietà comuni a strutture matematiche disparate come il gruppo additivo dei reali, il gruppo degli interi modulo un intero m , il gruppo delle isometrie del piano, il gruppo delle permutazioni di un insieme in sé, e così via. Spesso i teoremi sono estensioni e generalizzazioni di teoremi già noti ma che si riferivano casi particolari. In tale senso si produce nuova matematica e nuovi e più generali teoremi.

Intuizione.

Euclide. In Euclide l'intuizione è lo strumento mediante il quale si perviene alla conoscenza degli enti matematici e delle loro proprietà fondamentali. Che gli assiomi, per meglio dire i postulati, siano fortemente intuitivi è cosa essenziale perché solo questa loro immediatezza assicura la bontà di tutta la teoria che si vuole costruire (e rende possibile "postulare" per una loro accettazione). In definitiva la scelta del sistema di assiomi poggia direttamente sulla intuizione ed è, in un certo senso, ad essa successiva. Naturalmente l'intuizione non deve invece essere uno strumento di dimostrazione anche se essa conserva un ruolo euristico (in ciò Euclide non differisce dai moderni matematici).

Comunque l'intuizione di un buon matematico non conduce ad errori ed è solo per motivi di "purezza intellettuale" che non può servirsene in una trattazione scientifica di un argomento.

Hilbert. In Hilbert, poiché il metodo assiomatico ha lo scopo di fondare un settore della matematica preesistente alla assiomatizzazione, un adeguato sistema di assiomi per la geometria deve solo essere abbastanza potente da permettere di ottenere tutti i teoremi della geometria euclidea. Non è detto perciò che i singoli assiomi debbano essere "intuitivi". Naturalmente, nel caso in questione, gli assiomi di Hilbert ereditano il carattere intuitivo degli assiomi di Euclide ma, se ne avesse avuto bisogno, Hilbert non avrebbe esitato ad introdurre un sistema di assiomi completamente diverso il cui requisito avrebbe dovuto essere solo di essere capace di dimostrare in maniera elegante e rapida l'insieme delle proposizioni della geometria euclidea. L'intuizione comun-

che può essere forviante perché può portare a dare per scontato qualche cosa che invece deve essere dimostrato.

Strutturalisti. In maniera ancora più radicale per la scuola strutturalista l'intuizione è cosa più da combattere che da utilizzare ed essa non ha assolutamente valore euristico. Ad esempio se ogni volta che si parla dei gruppi ci si immagina il gruppo additivo dei reali, allora si è portati ad attribuire ai gruppi più proprietà di quante siano deducibili dal dato sistema di assiomi (ad esempio la commutatività). In altri termini, poiché per gli strutturalisti una teoria deve descrivere una intera classe di strutture diverse tra loro, è forviante nelle dimostrazioni guardare troppo ad una particolare struttura.

Categoricità.

Euclide. In Euclide la categoricità del sistema di assiomi è una questione che, in un certo senso, non si pone. Infatti il suo sistema di assiomi aveva l'unico scopo di descrivere in modo adeguato il modello fornito dall'idea platonica di punto, retta, circonferenza. Tale modello è il punto di partenza ed era considerato l'unico possibile. Pertanto non poteva mai porsi il problema della esistenza di un modello degli assiomi di Euclide diverso da quello naturale (non esistevano due mondi delle idee!).

Hilbert. Hilbert invece si pone esplicitamente la questione della categoricità, anzi la considera un requisito essenziale per il suo sistema in quanto la bontà del sistema di assiomi proposto consiste proprio nella capacità di individuare esattamente l'oggetto matematico a cui si vuole dare una base solida, ad esempio per quanto riguarda la geometria individuare il piano euclideo. Era anche importante provare la categoricità per il sistema di assiomi dei numeri interi, dei numeri reali ed altro.

Strutturalisti. In questo caso la bontà di tale teoria si misura nella sua capacità ad abbracciare quanti più "oggetti matematici". Un teorema della teoria dei gruppi è importante poiché fornisce informazioni su una vasta classe di strutture. Per gli strutturalisti se una teoria è categorica non è molto utile. Infatti se si vuole appurare se una asserzione α è conseguenza o meno degli assiomi della teoria allora tanto vale di andare a vedere direttamente se α vale su tale modello. Consideriamo ad esempio l'intero si-

stema di assiomi della geometria di Hilbert; si sa che tale sistema è categorico e che quindi tutti i modelli sono isomorfi tra loro ed in particolare sono equivalenti al modello analitico (cioè ad R^2). Allora dovendo verificare se una proposizione è vera oppure no, non conviene cercare di dedurla dal sistema di assiomi; è più semplice verificarne direttamente la validità su R^2 mediante alcuni semplici calcoli.

Completezza

Euclide. Anche se non esplicitamente menzionata è evidente che nella tradizione Euclidea la completezza di una teoria matematica era considerata un requisito importante. Un sistema di assiomi per la geometria per cui esista una asserzione non dimostrabile e non confutabile è evidentemente inadeguato a rappresentare l'intuizione geometrica. D'altra parte da Euclide fino agli inizi dell'ottocento si era convinti che questo assioma fosse in realtà dimostrabile e gli sforzi per dimostrarlo mostravano che si aveva il convincimento della completezza della teoria ottenuta cancellandolo.

Hilbert. In Hilbert la completezza è essenziale in quanto è essenziale la categoricità che, appunto, comporta la completezza. Se una teoria non è completa non può certo sperare di fondare qualche struttura matematica.

Strutturalisti. Anche per gli strutturalisti può essere ripetuto quanto detto per la questione della categoricità. Una buona teoria non deve essere completa in quanto deve essere la parte comune di più teorie diverse tra loro.

Consistenza.

Euclide. Anche il problema della consistenza non poteva essere problema riguardante Euclide. Infatti il sistema di postulati proposto era un elenco di proprietà di un modello preesistente di cui l'uomo ha diretta ed innata conoscenza. Ma se un sistema di assiomi viene costruito isolando alcune proprietà di un dato modello allora esso è necessariamente consistente. Se infatti fossero deducibili sia una formula α che una sua negata $\neg\alpha$, allora per il modello di partenza esisterebbe una proposizione che è sia vera che falsa, il che non può accadere.

Hilbert. Per Hilbert il problema della consistenza è essenziale, non esiste un mondo geometrico da descrivere, esiste solo un sistema di assiomi e niente assicura che tale sistema sia stato scelto male e che ci si accorga che esso permette di dimostrare sia che valga una cosa che valga il contrario. Facciamo un esempio, e supponiamo che Hilbert avesse messo tra gli assiomi della sua teoria anche l'assioma

$\alpha = \text{"l'insieme dei punti nel piano è numerabile"}$.

La teoria risultante T sarebbe apparsa degna di interesse e si sarebbero prodotti molti teoremi di tale teoria. Purtroppo, Hilbert ad un certo punto sarebbe giunto a dimostrare (cosa che fa nel suo libro) che T permette di introdurre un sistema di coordinate che utilizza il campo dei numeri reali. Ciò avrebbe comportato la validità di

$\neg\alpha = \text{"l'insieme dei punti nel piano non è numerabile"}$

È ovvio allora che la teoria T , non potendo ammettere modelli, sarebbe risultata priva di ogni interesse.

D'altra parte non si può provare la consistenza di un sistema di assiomi facendo ricorso ad un modello perché non si può utilizzare la cosa da fondare (il modello) per giustificare un sistema di assiomi (la teoria) che ha come scopo proprio quello di fare a meno del modello.

Strutturalisti. Molto meno importante è la questione della consistenza per gli strutturalisti. È vero che tale questione in linea di principio si pone, ma, di fatto, una teoria viene proposta solo allo scopo di unificare lo studio di una serie di modelli matematici preesistenti. Pertanto tale teoria sarà, come per Euclide, automaticamente consistente. Ad esempio, il problema della consistenza della teoria dei gruppi non si pone in quanto tale teoria nasce successivamente alla considerazione di alcuni gruppi concreti.

Il problema dei fondamenti.

Sia in Euclide che in Hilbert l'assiomatizzazione risponde al problema di dare un fondamento sicuro all'intero edificio della matematica. Negli strutturalisti l'atteggiamento è più pragmatico e "locale". Essi non pretendono di dare una volta per tutti un sistema di assiomi su cui fondare tutta la matematica. Piuttosto essi pensano che la matematica si possa smembrare in diversi settori i quali si possono esaminare indipendentemente uno dall'al-

tro. Ad esempio J. Dieudonné, uno dei fondatori dello strutturalismo in matematica, afferma

Il matematico moderno si sente così perfettamente in pace con la sua coscienza e non si preoccupa affatto di tutti gli pseudo-problemi che hanno preoccupato i suoi predecessori.

Gli pseudo-problemi di cui si parla sono i paradossi che non sembrano interessare più di tanto gli strutturalisti. Se anche una singola teoria risultasse essere contraddittoria, ciò non comporterebbe il crollo di tutto l'edificio matematico, ma semplicemente un piccolo cambio di indirizzo (un indebolimento di qualche assioma). Inoltre mentre la dimostrazione della coerenza dei sistemi di assiomi per l'aritmetica e la geometria presentano difficoltà insormontabili, lo stesso non avviene per molte "piccole" teorie esaminate dagli strutturalisti. La loro coerenza si può mostrare spesso semplicemente esibendo dei modelli finiti (come avviene per la teoria dei gruppi, dei campi, degli anelli ed altro). E poi le questioni relative ai fondamenti della matematica non devono infastidire i matematici di professione. Si provvederà ad addestrare appositamente del personale "paramatematico" (i logici) che se ne occupino.

Tutte le questioni come la non contraddizione delle teorie . . . ed in generale tutto ciò che concerne la teoria della dimostrazione fanno parte ora di una scienza completamente separata dalla matematica, la metamatematica; questa nuova disciplina non cessa di svilupparsi ed ha già fornito numerosi risultati nuovi e pieni di interesse; ma è perfettamente lecito al matematico ignorarla completamente senza essere per nulla preoccupato nelle sue ricerche (è sempre Dieudonné che parla).

Nota. La distinzione tra approccio fondazionale ed approccio strutturalista al metodo assiomatico è di un certo interesse da un punto di vista didattico. Infatti il metodo assiomatico nelle scuole può essere introdotto o tramite un sistema di assiomi per la geometria euclidea (Hilbert) oppure tramite la nozione di gruppo o di insieme ordinato (Bourbakismo). Nel primo caso la scelta fondazionale presenta difficoltà se gli studenti non sono stati stimolati verso interessi epistemologici. Infatti per essi sembrerà strano che si debba perdere tanto tempo per dimostrare cose che appaiono ovvie. Inoltre l'esistenza di un unico modello renderà

difficile far capire questioni di indipendenza. Infine il carattere estremamente intuitivo di tale modello farà scomparire in modo completo la questione dell'indipendenza. Nel secondo caso l'approccio unificante del metodo assiomatico in chiave strutturalista può essere compreso dagli alunni solo se prima sono stati esibiti molti esempi di strutture diverse che però sono suscettibili di una trattazione unificata. Ad esempio dovrebbero essere prima introdotti alcuni gruppi particolari come quello delle simmetrie di una figura, quello delle permutazioni di un insieme, quello dei movimenti del piano e così via.

CAPITOLO 6

LA MATEMATICA COME CALCOLO CON PAROLE¹

“se si lodano gli uomini che hanno determinato il numero di corpi regolari, che non ha utilità alcuna se non in quanto è piacevole a contemplarsi, quanto sarà più meritorio ridurre a leggi matematiche il ragionamento umano, che è ciò che di più eccellente e di più utile possediamo.” W. G. Leibniz.

1. Hilbert contro l'infinito

Abbiamo visto che il problema della non contraddittorietà di un dato sistema di assiomi T è centrale per la concezione ipotetico-deduttiva di Hilbert. Ora i problemi di non contraddittorietà di una teoria si risolvono usualmente mostrando un modello di tale teoria, mostrando cioè che essa è soddisfacibile.² Ad esempio la consistenza della geometria non euclidea viene provata dal modello di Klein o da uno dei tanti modelli di cui abbiamo già parlato.³ D'altra parte tali modelli vengono costruiti a partire dal modello euclideo e quindi la dimostrazione di non contraddittorietà delle geometrie non euclidee è valida a patto che la geometria euclidea sia affidabile. In altre parole:

riusciamo a provare che la teoria delle geometrie non euclidee è consistente solo provando che il sistema di assiomi per la geometria euclidea è consistente.

A questo punto si potrebbe dire che è facile provare tale consistenza, basta costruire il modello analitico basato sui numeri reali. Ma come giustificare i numeri reali? Coerentemente con il metodo assiomatico, possiamo fornire un sistema di assiomi per i nu-

¹ Questo capitolo riguarda la logica matematica, un argomento i cui contenuti non possono essere certamente concentrati in poche pagine. Pertanto ci si limiterà solo a fornire alcune delle idee-base e ciò comporta che si scriveranno solo definizioni e gli enunciati dei principali teoremi. Per chi sia interessato si suggerisce uno dei tanti libri in commercio.

² Come abbiamo già detto nel capitolo precedente, una teoria si dice *soddisfacibile* se ammette un modello, si dice *consistente* se da tale teoria non è possibile ricavare una contraddizione. Ovviamente se una teoria è soddisfacibile è anche consistente e pertanto per provare la consistenza di una teoria è sufficiente esibirne un modello.

³ Cioè del sistema di assiomi che contiene la negazione del quinto postulato da non confondere con il sistema di assiomi che non contiene tale postulato.

meri reali (ad esempio tramite la nozione di campo completo). Ovviamente ciò conduce al seguente fatto:

riusciamo a provare che il sistema di assiomi della geometria euclidea è consistente solo provando che la teoria dei campi completi è consistente.

Ancora una volta sembra che non ci siano difficoltà: infatti abbiamo mostrato come, ad esempio con il metodo delle sezioni, sia possibile il campo dei reali e quindi un modello della teoria dei campi completi. Purtroppo però sia per dimostrare l'esistenza di una terna di Peano sia per effettuare una delle tante costruzioni dei numeri reali dobbiamo servirci della teoria degli insiemi. In definitiva, se ci potessimo fidare della teoria degli insiemi, avremmo risolto tutti i problemi e ci potremmo fermare:

riusciamo a provare la consistenza delle varie teorie utilizzate in matematica solo se la teoria degli insiemi è consistente.

Sfortunatamente la scoperta dei paradossi mostra che la nozione di insieme è alquanto inaffidabile e che quindi si deve procedere ad una sua buona assiomatizzazione. Ora è vero che esistono diverse teorie assiomatiche degli insiemi che permettono di evitare tutti i paradossi fino ad ora noti, ma chi ci assicura che un giorno non vengano scoperti paradossi anche per tali nuove teorie? Il problema appare senza soluzioni.

Ora Hilbert pensò che tutte le difficoltà nascessero dalla considerazione dell'infinito attuale che già tanta diffidenza aveva suscitato da Pitagora in poi. Infatti, come abbiamo già osservato, nessuno dubita della consistenza della teoria dei gruppi perché non è difficile fornire "concretamente" esempi di gruppi finiti. Esaminiamo in proposito un passo dall'articolo di Hilbert *Sull'infinito* apparso nel 1925.⁴

⁴ In questo passo viene fatto riferimento alla “operazioni sull'infinitamente piccolo” tipiche dell'analisi matematica. Si deve pensare che prima dell'attuale definizione di limite (in termini di quantificatori universali ed esistenziali) le nozioni di limite e di derivata venivano viste come il risultato di operazioni fatte su infiniti o infinitesimi, i quali, come poi sarà fatto in modo rigoroso dall'analisi non standard, venivano visti come particolari tipi di quantità.

“Proprio come le operazioni sull'infinitamente piccolo sono state sostituite da operazioni sul finito che danno luogo esattamente agli stessi risultati e alle stesse eleganti relazioni formali, così in generale i metodi deduttivi basati sull'infinito devono essere sostituiti con procedimenti finiti che diano gli stessi risultati, che cioè rendano possibili le stesse catene di dimostrazioni e gli stessi metodi per ottenere formule e teoremi.

Questo è lo scopo della mia teoria, essa si propone di dare definitivamente sicurezza al metodo matematico . . .”

Hilbert vede nella chiarificazione del concetto di infinito una questione fondamentale il cui interesse non è solo matematico.

“Le considerazioni precedenti intendono solo affermare che la chiarificazione definitiva della natura dell'infinito non riguarda esclusivamente l'ambito degli interessi scientifici specializzati ma è necessaria per la dignità stessa dell'intelletto umano.”

“D'altra parte l'esistenza dei paradossi della teoria degli insiemi sembra spingere al rifiuto dell'infinito attuale. C'è tuttavia un modo soddisfacente per evitare i paradossi senza tradire la nostra scienza. Il punto di vista utile per la scoperta di tale modo ed il desiderio che ci mostra la via da prendere sono:

1. Se c'è la più piccola speranza, esamineremo accuratamente tutte le definizioni e i metodi deduttivi fecondi, li curemo, li potenzieremo e li renderemo utili. Nessuno potrà cacciarci dal paradiso che Cantor ha creato per noi.⁵

2. Dobbiamo estendere a tutta la matematica quella sicurezza dei metodi dimostrativi che è propria della teoria elementare dei numeri, di cui nessuno dubita e in cui contraddizioni e paradossi sorgono solo per negligenza.”

⁵ Ho sottolineato io questa frase poiché evidenza che il punto di vista di Hilbert è completamente differente da quello degli intuizionisti che invece pensavano si dovesse cancellare dalla matematica tutto ciò che va oltre il numerabile ed il costruttivo.

Il punto di vista di Hilbert è chiaro. Per prima cosa il rigore in matematica non si deve ottenere semplicemente eliminando quella parte della matematica e quei metodi che, pur essendosi rivelati fecondi, non risultano avere basi sicure. Pertanto, nonostante i paradossi, *“Nessuno potrà cacciarci dal paradiso che Cantor ha creato per noi”*. L'atteggiamento di Hilbert è pragmatico, se certi metodi si sono rivelati utili allora devono essere accettati.

“In effetti il successo è essenziale perché, in matematica come altrove, esso costituisce la corte suprema di fronte a cui tutti si inchinano.”

D'altro lato è indiscutibile che il rigore e la sicurezza si possono ottenere solo facendo riferimento ai metodi finitari propri dei numeri interi. Come fare per conciliare le due cose apparentemente contraddittorie?

2. L'infinito è solo una parola

Abbiamo già visto una citazione di Hilbert che si riferisce alla nozione di limite dichiarando che in tale nozione il coinvolgimento dell'infinito è solo aparente. Infatti se si scrive $\lim_{n \rightarrow \infty} 1/n = 0$ non si deve intendere che quando n raggiunge l'infinito allora il valore di $1/n$ raggiunge 0. Piuttosto tale uguaglianza è solo un modo breve per indicare che:

per ogni $\varepsilon > 0$ esiste m tale che $1/n \leq \varepsilon$ per ogni $n \geq m$.

E' evidente allora che in tale caso il simbolo ∞ non denota niente ma è solo un aiuto linguistico per rappresentare una situazione in cui l'infinito non compare in nessun modo.

Un discorso simile, come osserva ancora Hilbert, può essere fatto in relazione alla nozione di "punto all'infinito" elaborata dalla geometria proiettiva. Come è noto, il piano proiettivo viene definito aggiungendo all'insieme dei punti del piano euclideo dei punti ideali, detti "punti all'infinito". In tale modo invece di dire che due rette sono parallele diciamo che hanno in comune un punto all'infinito. Tecnicamente ciò si ottiene al modo seguente.

Definizione 2.1. Chiamiamo *punto all'infinito* del piano euclideo ogni fascio completo di rette parallele. Diciamo che un punto all'infinito P appartiene ad una retta r se r appartiene al fascio P .

L'introduzione dei punti all'infinito permette di semplificare e rendere simmetrici gli assiomi della geometria. Ad esempio due rette si incontrano sempre in un punto (che è finito se le rette non sono parallele ed infinito se le rette sono parallele).⁶ Anche in questo caso non si pretende che i punti all'infinito siano realmente esistenti, essi sono strumenti linguistici, enti ideali la cui introduzione è utile per ottenere risultati e per avere una trattazione più efficace della geometria.⁷

La proposta di Hilbert è pertanto di considerare l'infinito un ente ideale, per meglio dire un oggetto linguistico da manipolare secondo certe regole. In altre parole si tratta di spostare il ruolo del linguaggio il quale, da strumento di indagine del mondo degli enti matematici, deve diventare esso stesso oggetto di investigazione. Oggetto di studio dovranno essere i "segni concreti" che rimangono comunque oggetti finiti da maneggiare con un numero finito di regole.

“Questa è la filosofia che ritengo necessaria non solo per la matematica ma per ogni pensiero, per ogni comprensione e per ogni comunicazione che rientrano nell'ambito della scienza. In base ad essa, in particolare, oggetto della nostra considerazione matematica sono gli stessi segni concreti la cui forma, in virtù del nostro approccio, è immediatamente chiara e riconoscibile.”

Fino a qui non esistono, come lo stesso Hilbert sottolinea, grandi differenze con la tradizione algebrica. Anche la semplice risoluzione di una equazione di primo grado consiste in una manipolazione di equazioni (oggetti linguistici) secondo certe regole che permettono di passare da una equazione ad un'altra. Lo stesso si può dire anche di un semplice calcolo. La grossa novità nasce dal fatto che per "segni concreti" Hilbert intendeva non solo equazioni ma anche espressioni linguistiche molto più complicate che coinvolgevano i quantificatori, i connettivi logici (ad esempio la negazione, la congiunzione, la disgiunzione). Inoltre tra le regole di manipolazione linguistica non considerava solo quelle di carattere algebrico (quali ad esempio la proprietà associativa, la proprietà distributiva ...) ma anche quelle di natura logica quali le re-

⁶ Inoltre in tale modo si ottiene un potente ed elegante strumento per la trattazione delle curve algebriche.

⁷ Discorso analogo vale per l'introduzione dell'unità immaginaria.

gole di inferenza. Arriviamo pertanto al punto fondamentale: il calcolo logico.

“Certo questo fu sviluppato in origine per motivi del tutto differenti. I suoi segni furono introdotti originariamente solo per scopi di comunicazione. Tuttavia è coerente col nostro punto di vista non attribuire alcun significato ai segni logici, così come non se ne è attribuito alcuno ai segni matematici, e dichiarare che anche le formula del calcolo logico sono elementi ideali che di per sé non significano niente. Col calcolo logico abbiamo un linguaggio simbolico che permette di tradurre in formule le asserzioni matematiche e di esprimere le deduzioni logiche mediante processi formali.”

In definitiva non bastava ridurre la matematica a linguaggio, ma era anche necessario formalizzare i processi deduttivi che permettevano la manipolazione di tale linguaggio. Più precisamente una particolare teoria matematica veniva vista come:

- un insieme finito di espressioni linguistiche (gli assiomi propri della teoria)
- un insieme fissato di espressioni (gli assiomi logici)
- il tutto da manipolare tramite determinate regole (le regole di inferenza).

Tale apparato permette di produrre altre espressioni (i teoremi). In tale modo qualunque teoria (anche quelle che parlano di oggetti infiniti) diviene un oggetto finito e quindi passibile di essere esaminato nella sua interezza. Il problema della consistenza diviene allora quello di esaminare tale oggetto finito e vedere se tra le sue proprietà vi è anche quella della consistenza. Un programma questo che sembra abbastanza ragionevole.

3. Nuovi oggetti matematici: parole e linguaggi

Abbiamo visto che per Hilbert la matematica, in particolare la dimostrazione matematica, può essere vista come un particolare sistema che agisce sulle parole di un linguaggio. Ora la nozione di “parola” e di “linguaggio” non sembrano appartenere all’universo matematico e quindi si pone il problema di come possa essere reso rigoroso il discorso di Hilbert. E’ necessario a tale scopo elaborare una teoria matematica dei linguaggi e la cosa può essere fatta al modo seguente. Si parte da un “alfabeto”, cioè

un insieme finito di simboli. Ad esempio nella lingua italiana un alfabeto è costituito dalle lettere $a, b, c, \dots$. In lingua greca è costituito dalle lettere $\alpha, \beta, \gamma, \dots$. Un esempio di alfabeto “piccolo” è l’alfabeto Morse $A = \{., -\}$ costituito da un punto e da una linea che una volta veniva usato per il telegrafo. In definitiva possiamo chiamare alfabeto un qualunque insieme finito. Dato un alfabeto possiamo poi costruire parole. Ad esempio con quello della lingua italiana possiamo scrivere parole come “cane”, “canne”, “cena”, “acne”, “acccanne”.. (quest’ultima è una parola senza significato). Dovendo definire la nozione di parola in termini insiemistici osserviamo che due parole come “cane” e “cena”, che pur avendo le stesse lettere le hanno in posizione diversa, devono essere considerate diverse. Sono pure da considerare diverse due parole come “cane” e “canne” in cui le lettere coincidono anche nell’ordine ma appaiono un numero diverso di volte. Ciò mostra che non è certamente possibile chiamare parola un insieme finito di elementi dell’alfabeto. Infatti, poichè due insiemi con gli stessi elementi sono uguali, tutti gli insiemi $\{c, a, n, e\}$, $\{c, a, n, n, e\}$, $\{c, e, n, a\}$, $\{a, c, n, e\}$, $\{a, c, c, a, n, n, e\}$ coincidono. Allora, per potere rappresentare la nozione di parola conviene usare la nozione di n -pla che permette appunto di tenere conto sia dell’ordine che dell’eventuali ripetizioni. Si arriva quindi alla seguente definizione.

Definizione 3.1. Sia A un insieme finito (che chiamiamo *alfabeto*). Allora gli elementi di A^n vengono detti *parole di lunghezza n* nell’alfabeto A . Indichiamo con A^+ l’insieme $\bigcup_{n \in \mathbb{N}} A^n$ di tutte le possibili parole.

Una volta che abbiamo soddisfatto la mania dei matematici di ridurre tutto alla teoria degli insiemi, continueremo a rappresentare le parole al modo solito, scrivendo una parola come $(a_1, a_2, \dots, a_n)$ con $a_1, a_2, \dots, a_n$.⁸ Ad esempio per indicare la parola (c, a, n, e) scriveremo semplicemente *cane*.

⁸ Il fatto che “rendere rigoroso” significhi ridurre alla teoria degli insiemi è una credenza, un po’ ridicola, dei matematici. Assumere come primitiva direttamente la nozione di parola sarebbe altrettanto o forse maggiormente rigoroso viste le basi incerte della teoria degli insiemi. E’ quello che usualmente viene fatto nei linguaggi di programmazione dove viene assunta come primitiva la nozione di “parola” (equivalente-

Possiamo immaginare l'alfabeto A come l'insieme dei caratteri che sono sulla tastiera di un computer ed una parola come la sequenza di lettere che compaiono sullo schermo quando si batte sulla tastiera. In particolare è possibile anche considerare il “simbolo spazio-vuoto” che corrisponde al lasciare uno spazio vuoto tra due lettere. In questo caso si chiama parola anche quella che è per noi una frase, cioè una sequenza di parole separate da spazi vuoti. Allora possiamo far rientrare nel termine “parola” anche frasi del tipo

“il cane corre”, “enac corre il”, ...

Ad esempio il contenuto di questo libro può essere visto come una unica parola nell'alfabeto che si ottiene usando i soliti caratteri della tastiera più un “carattere” per denotare lo spazio vuoto.

Un linguaggio non si caratterizza solo da un alfabeto ma anche da un criterio con cui si stabilisce se una parola o frase è da considerare accettabile. Stabilito tale criterio, allora un linguaggio si definisce come l'insieme delle parole (se si vuole di frasi) che sono considerate accettabili.

Definizione 3.2. Dato un alfabeto A chiamiamo *linguaggio formale su A* un sottoinsieme $\mathcal{L}$ di A^+ .

Se $\mathcal{L}$ contiene solo un numero finito di parole, allora un criterio di accettabilità si può ottenere direttamente tramite una lista delle parole che sono in $\mathcal{L}$. Ad esempio nel caso della lingua italiana l'insieme corretto delle parole coincide con l'insieme delle parole presenti in un buon vocabolario. Nel caso il linguaggio contiene un insieme infinito di parole (come in generale avviene quando si considerano frasi), allora è un po' più complicato stabilire un criterio di accettabilità ed in generale ci si riferisce alla nozione di *grammatica* che però non considereremo in questo libro. Esaminiamo ora una ovvia proprietà dei linguaggi formali.

Proposizione 3.3. L'insieme A^+ delle parole nell'alfabeto A è numerabile. Pertanto ogni linguaggio formale è finito o numerabile.

mente, il tipo “stringa”, “lista”, ...). Dovendo poi introdurre il tipo “insieme” lo si fa derivare da quello di “parola” introducendo una opportuna relazione di equivalenza tra parole (o, equivalentemente, di riduzione a forma normale).

Dim. Sia $a \in A$, allora la corrispondenza f che associa ad ogni $n \in N$ la parola $aaaa..a$ che si ottiene ripetendo n volte la lettera a è una funzione iniettiva di N in A^+ . Quindi A^+ ha potenza maggiore o uguale del numerabile. Vogliamo provare che A^+ ha anche potenza minore o uguale al numerabile, cioè che esiste una funzione iniettiva di A^+ in N . Esistono diversi modi per fare questo che spesso prendono il nome di *codifiche*. Un modo vicino allo stile dei matematici è considerare la successione $p_1, p_2, \dots$ dei numeri primi scritta in ordine crescente. Associamo poi in modo iniettivo ad ogni lettera in A un numero intero, cioè consideriamo una funzione iniettiva $c : A \rightarrow N$. Essendo A finito non ci sono difficoltà a trovare una tale funzione. Successivamente associamo ad ogni parola $a_1 \dots a_n$ il numero $p_1^{c(a_1)} \cdot \dots \cdot p_n^{c(a_n)}$. Tale corrispondenza è iniettiva (anche se non suriettiva). Ad esempio se poniamo $c(a) = 1, c(b) = 2, c(c) = 3, \dots$ allora alla parola “*bacca*” viene associato il numero $2^2 \cdot 3^1 \cdot 5^3 \cdot 7^3 \cdot 11^1$.

Un modo più “informatico” ed intuitivo di procedere è il seguente. Definiamo una funzione iniettiva $c : A \rightarrow \{0,1\}^+$ che associ ad ogni elemento $a \in A$ una parola nell’alfabeto $\{0,1\}$ e supponiamo che tutte le parole $c(a)$ abbiano la stessa lunghezza e che inizino con 1. Consideriamo la funzione che associa ad ogni parola $a_1 a_2 \dots a_n$ in A^+ la parola $c(a_1)c(a_2)\dots c(a_n)$ in $\{0,1\}^+$. Infine interpretiamo tale parola come un numero scritto in base 2 (ma il discorso funziona anche se la si interpreta in base 10). In tale modo abbiamo definito una funzione iniettiva di A^+ in N .⁹ Ad esempio supponiamo, sempre riferendoci all’alfabeto della lingua italiana che

$$c(a) = 111111, c(b) = 111011, c(c) = 100011, \dots$$

⁹La tecnica che ora abbiamo utilizzato per dimostrare la proposizione corrisponde al modo concreto di funzionamento dei computer. Infatti quando si preme un tasto della tastiera l’impulso inviato al computer è proprio una sequenza di 0 ed 1. Dopo avere scritto un documento (un romanzo una poesia od altro), nella memoria del computer è immagazzinata una serie lunghissima di 0 e di 1 corrispondente alla sequenza di lettere scritte (ma anche agli spazi vuoti). Più precisamente nei sistemi di scrittura (come il Word) esistono anche caratteri che vengono chiamati “di controllo” e che servono a modificare i caratteri che compaiono sullo schermo. Ad esempio esistono caratteri che creano il corsivo, il grassetto e così via.

In questo modo associamo alla parola “bacca” il numero, scritto in base 2,

111011111111100011100011111111.

□

Si osservi che il procedimento con cui abbiamo associato ad ogni parola un numero naturale è “effettivo”, nel senso che corrisponde ad un preciso algoritmo che un qualunque calcolatore sarebbe capace di eseguire. In questo caso si usa parlare di “codifica” di un linguaggio. Si usa anche considerare una “decodifica” che associa ad ogni numero intero m l’elemento del linguaggio formale $\mathcal{L}$ ottenuto al modo seguente (ci riferiamo al primo tipo di codifica sopra esposto)

- si fissa una parola p in $\mathcal{L}$
- si effettua una scomposizione $m = p(1)^{i(1)} \dots p(n)^{i(n)}$ di m in prodotto di successivi numeri primi
- se $i(1), \dots, i(n)$ sono codici di lettere in A (cioè se esistono $a_1, \dots, a_n$ in A tali che $g(a(1)) = i(1), \dots, i(a(n)) = i(n)$) e se la parola $a_1 \dots a_n$ è una parola in $\mathcal{L}$, allora si assume tale parola come decodifica di m
- altrimenti assumiamo per convenzione che la decodifica di m sia p .

Da notare che una tale decodifica è un processo effettivo solo se esiste un metodo per decidere se una parola appartiene ad $\mathcal{L}$ o meno ($\mathcal{L}$ è decidibile).

Corollario 3.4. L’insieme dei possibili romanzi è numerabile. L’insieme delle possibili poesie è numerabile. L’insieme dei possibili programmi di un linguaggio di programmazione è numerabile.

4. Rappresentabilità, definibilità e numerabilità

Possiamo utilizzare la proposizione 3.3 per provare in modo immediato che alcuni insiemi sono numerabili. A tale scopo ci serviamo della nozione di “rappresentabile” o quella analoga di “definibile” in un alfabeto A . Lasciamo il significato di tali nozioni a livello intuitivo. Quello che interessa è che, dato un insieme X ed un alfabeto A , una rappresentazione (o definizione) di un elemento di X in A sia una parola su A . Inoltre supponiamo che elementi diversi di X abbiano rappresentazioni (definizioni) diverse. Per fare un esempio, riferendoci alla rappresentazione decimale dei

numeri, se si accetta un alfabeto A con i simboli 0, 1, 2, 3, 4, 5, 6, 7, 8, 9 allora ogni naturale è rappresentabile in A . Se aggiungiamo ad A i simboli $+$ e $-$, allora otteniamo un alfabeto in cui ogni numero intero relativo è rappresentabile. Se all'alfabeto aggiungiamo anche il simbolo $/$, allora anche i razionali sono rappresentabili. Esempi di definizioni sono i seguenti. Il numero 3 può essere definito come “quel numero positivo il cui quadrato è nove”. Il numero $+\sqrt{2}$ può essere definito come “quel numero positivo il cui quadrato è 2”.

Teorema 4.1. Sia X un insieme i cui elementi siano rappresentabili (definibili) in un linguaggio formale. Allora X è enumerabile (cioè finito o numerabile).

Dim. E' possibile considerare una funzione che associa ad ogni elemento in X una delle sue rappresentazioni (definizioni) in A^+ . Si ottiene una funzione iniettiva di X in A^+ . Poiché A^+ è un insieme numerabile, questo prova che X è finito o numerabile. $\square$

Applicando tale criterio possiamo dimostrare la seguente proposizione.

Proposizione 4.2. L'insieme $P_f(N)$ delle parti finite di N è numerabile.

Dim. Ogni insieme finito di numeri interi può essere rappresentato tramite una parola del tipo $\{n_1, \dots, n_p\}$ con $n_1 < n_2 < \dots < n_p$ cioè tramite una parola nell'alfabeto costituito dalle due parentesi $\{, \}$, dalla virgola e dalle cifre 0, 1, 2, 3, 4, 5, 6, 7, 8, 9. Pertanto $P_f(N)$ è enumerabile. Poiché è ovvio che $P_f(N)$ non è finito, $P_f(N)$ è numerabile.¹⁰ $\square$

E' possibile anche ridimostrare in modo più semplice alcune delle proposizioni provate nel paragrafo precedente.

¹⁰ Concretamente, dato un insieme finito lo “descrivo” al computer con una parola del tipo $\{n_1, \dots, n_p\}$ con $n_1 < n_2 < \dots < n_p$. Tale parola viene immagazzinata nei registri di memoria sotto forma di una sequenza di 0 ed 1, sequenza che possiamo interpretare come un numero. Ciò fornisce una corrispondenza di $P_f(N)$ in N .

Proposizione 4.3. L'insieme Z dei numeri relativi è numerabile. L'insieme Q dei razionali è numerabile. L'insieme delle n -ple di numeri razionali è numerabile.

Dim. Già abbiamo osservato che gli elementi di Z e quelli di Q sono rappresentabili in un opportuno alfabeto. Per rappresentare l'insieme di tutte le possibili n -ple di numeri razionali è sufficiente aggiungere all'alfabeto le parentesi $(,)$ e la virgola.

Ricordiamo che un numero reale si dice *algebrico* se è radice di un polinomio a coefficienti razionali. Un numero è *trascendente* se non è algebrico. Poiché ogni polinomio a coefficienti razionali si può ridurre ad un polinomio equivalente a coefficienti interi, possiamo definire un numero algebrico anche come un numero che sia radice di un polinomio a coefficienti interi. Ad esempio $\sqrt{5}$ è un numero algebrico perché è soluzione dell'equazione $x^2 - 5 = 0$. E' algebrico anche un numero razionale p/q qualunque perché è soluzione dell'equazione $q \cdot x - p = 0$.

Proposizione 4.4. L'insieme Al dei numeri reali algebrici è numerabile. Pertanto l'insieme dei numeri trascendenti ha la potenza del continuo.

Dim. Ogni numero algebrico può essere definito da una proprietà del tipo “la terza radice del polinomio $p(x)$ ”. L'insieme di tali definizioni è un insieme di parole in un opportuno alfabeto. Ciò prova che Al è enumerabile. Poiché Al contiene tutti i razionali siamo sicuri che Al è numerabile. $\square$

Possiamo provare anche le seguenti sorprendenti conseguenze, in ambito informatico, dei teoremi ora dimostrati. Dato un linguaggio di programmazione, chiamiamo *decidibile* un sottoinsieme X di N per il quale esista un programma in grado di dirci se n appartiene ad X oppure no. Possiamo vedere un insieme non decidibile come un insieme X talmente complicato che non potrà mai essere scritto un programma capace di dirci se un dato elemento appartiene ad X oppure no. Si pone allora il problema se tali insiemi “patologici” esistano oppure no. Il seguente teorema fornisce una risposta.

Teorema 4.5. Il numero di insiemi non decidibili è maggiore di quello degli insiemi decidibili.

Dim. Un programma non è altro che una parola in un alfabeto finito e quindi ogni insieme decidibile può essere rappresentato da una parola su tale alfabeto (ovviamente insiemi decidibili diversi hanno programmi diversi). Da ciò segue che la classe dei sottoinsiemi decidibili è enumerabile. Il complemento di tale classe non può essere numerabile poiché altrimenti $P(N)$ sarebbe numerabile in contrasto con il teorema di Cantor. In definitiva la classe degli insiemi non decidibili ha potenza maggiore di quella degli insiemi decidibili. $\square$

Chiamiamo *computabile* una funzione di N in N per la quale esista un programma capace di computarla. Procedendo come per il teorema ora dimostrato, possiamo dimostrare anche il seguente teorema.

Teorema 4.6. Il numero delle funzioni non computabili è maggiore del numero delle funzioni computabili.

5. Linguaggio ed apparato deduttivo per la logica formale

Riferendoci alla matematica, un linguaggio formale che vada bene dovrà contenere

- nomi per oggetti matematici (come $3, 13, \pi, \emptyset, e$),
- nomi per funzioni ed operazioni (come $\log, +, \cdot, \text{sen}, x_2$)¹¹
- nomi per relazioni (come $=, \leq, \geq, \supseteq$).

Sembra infatti ragionevole poter scrivere asserzioni come " $\log(2+3) \geq 0$ " oppure " $\log(2+3) = 2 \cdot 3$ " che si avvalgono del nome di funzione $\log$, dei nomi di numeri $2, 3$ e 0 , di un simbolo $+$ per un'operazione binaria e dei simboli $\geq$ e $=$ per relazioni binarie. Ancora, in matematica sono di uso frequente le variabili ed i quantificatori (esiste, per ogni) che permettono ad esempio di scrivere proposizioni del tipo "esiste una soluzione dell'equazione $x^2-1=0$ ", in breve " $\exists x(x^2-1=0)$ ". Infine sono usati connettivi logici come "non", "e", "oppure", "implica" che consentono di costruire

¹¹ Non abbiamo escluso che in un alfabeto ci possano essere, come elementi, parole di un altro linguaggio come ad esempio, $\log, \text{sen}, \dots$. Infatti l'unica cosa che abbiamo richiesto è che un alfabeto sia un insieme finito.

asserzioni (composte) a partire da altre asserzioni. Ciò suggerisce le seguenti definizioni.

Definizione 5.1. Chiamiamo *alfabeto di un linguaggio del primo ordine* un alfabeto A costituito da:

- un insieme di simboli per denotare *variabili*, ad esempio $x, y, z, \dots$
- i *connettivi proposizionali* $\{\wedge, \vee, \rightarrow, \sim\}$
- il quantificatore esistenziale $\exists$
- la parentesi aperta (e la parentesi chiusa)

In più

- un insieme finito C di elementi detti *costanti*
- un insieme finito O di *nomi di operazione*
- un insieme finito P di *nomi di predicati*
- una *funzione arità* $ar: O \cup P \rightarrow \mathbb{N}$

Se l'arità di un nome di funzione è 2 allora è inteso che tale nome denota una funzione binaria e si preferisce parlare di *operazione binaria*. Se l'arità di un nome di predicato è 1 allora si dice che denota un predicato *monadico* o una *proprietà*. Se invece l'arità è 2 si parla di *relazione binaria*. Da notare che esistono tanti alfabeti, e quindi tanti linguaggi del primo ordine, quanti sono i modi di specificare le costanti, i nomi delle relazioni e delle operazioni. Avremo ad esempio un linguaggio adeguato alla teoria dei gruppi, uno per la teoria degli anelli, uno per le strutture ordinate, e così via. Per poter definire il linguaggio del primo ordine corrispondente ad un dato alfabeto dobbiamo prima definire il linguaggio dei termini. Possiamo vedere un termine come descrizione di un algoritmo per il calcolo di una funzione (se il termine contiene variabili) o di un elemento (se il termine non contiene variabili).

Definizione 5.2. Dato un alfabeto del primo ordine chiamiamo *linguaggio dei termini* l'insieme delle parole che si ottengono applicando più volte le seguenti regole:

- a) ogni variabile o costante è un termine
- b) se f è il nome di una funzione n -aria e $t_1, \dots, t_n$ sono termini allora $f(t_1, \dots, t_n)$ è un termine (notazione *prefissa*)
- c) se $\otimes$ è il nome di una funzione binaria e t_1, t_2 sono termini allora $(t_1) \otimes (t_2)$ è un termine (notazione *infissa*).

In matematica usualmente si usano solo nomi di funzioni di arità 1 (dette *funzioni* o *operazioni* unarie) o 2 (dette *operazioni binarie*). In questo ultimo caso si preferisce la notazione infissa. Ad esempio l'espressione $(x) - (\log((y) + (x)))$ è un termine. Infatti

- poiché y è un termine ed x è un termine allora $(y) + (x)$ è un termine
- poiché $(y) + (x)$ è un termine allora $\log((y) + (x))$ è un termine
- essendo x un termine e $\log((y) + (x))$ un termine, $(x) - (\log(y + x))$ è un termine.¹²

Utilizzando i termini, possiamo ora definire la nozione di linguaggio del primo ordine.

Definizione 5.3. Un *linguaggio del primo ordine*¹³ è l'insieme $\mathcal{L}$ delle parole su di un alfabeto del primo ordine definito dalle seguenti regole di formazione:

- a) se r è un predicato n -ario e $t_1, \dots, t_n$ sono termini allora $r(t_1, \dots, t_n)$ è una formula (*formula atomica in notazione prefissa*)

¹² Naturalmente la realtà dei linguaggi matematici è più flessibile e diversificata. Ad esempio il termine che abbiamo costruito si indica più semplicemente con $x - \log(y + x)$ in quanto esistono semplici regole e convenzioni per l'eliminazione di parentesi inutili. Ad esempio scriviamo $3 \cdot 5 + 2$ al posto di $((3) \cdot (5)) + (2)$ come invece richiederebbe la definizione 4. Inoltre esistono anche notazioni "*post-fisse*" per le operazioni unarie, ad esempio per la funzione fattoriale $x!$ e notazioni "*esponenziali*" come la funzione inverso x^{-1} . Tuttavia ai fini del nostro discorso ci atterremo alla Definizione 2 anche se non l'applicheremo in modo rigoroso.

¹³ Prende il nome di *linguaggio del secondo ordine* un linguaggio in cui sia possibile quantificare sui sottoinsiemi del dominio e non solo sugli elementi del dominio. Ciò comporta l'uso di diversi tipi di simbolo per le variabili che si riferiscono a sottoinsiemi e variabili che si riferiscono agli elementi. Ad esempio se utilizzo lettere maiuscole nel primo caso e minuscole nel secondo caso, la nozione di buon ordinamento potrà essere espressa dalla formula

$$\forall X \exists m (m \in X \wedge \forall z (z \in X \rightarrow m \leq z)).$$

Anche l'assioma di completezza, che si esprime dicendo che ogni sottoinsieme inferiormente limitato ammette estremo inferiore si esprime quantificando sui sottoinsiemi. Un altro importante esempio è l'assioma per le terne di Peano che riguarda il principio di induzione. Come abbiamo già accennato, la logica che si riferisce ad un linguaggio del secondo ordine prende il nome di *logica del secondo ordine*. In questo volume non si parla di questo tipo di logica e ci si limita alla *logica del primo ordine* cioè la logica che tratta dei linguaggi del primo ordine.

- b) se r è il nome di un predicato binario e t_1, t_2 sono termini allora $t_1 r t_2 \in \mathcal{L}$ (formula atomica in notazione infissa);
 c) se α e $\beta \in \mathcal{L}$ allora $(\alpha) \wedge (\beta)$, $(\alpha) \vee (\beta)$, $(\alpha) \rightarrow (\beta)$ e $\neg(\alpha)$ appartengono ad $\mathcal{L}$;
 d) se x è una variabile ed $\alpha \in \mathcal{L}$ allora $\exists x(\alpha) \in \mathcal{L}$.

Gli elementi di $\mathcal{L}$ vengono chiamati *formule ben formate* o, più semplicemente, *formule*.

I seguenti sono alcuni esempi di linguaggi del primo ordine utilizzati in matematica.

Esempi. Linguaggio usato per le strutture ordinate. È un linguaggio che contiene i soli simboli $\leq$ e $=$ di relazioni binarie. A volte si aggiunge anche una costante 0 (da interpretare come minimo elemento) e una costante 1 (da interpretare come massimo elemento) e pertanto $C = \{0, 1\}$. Poiché non ci sono nomi di operazioni, gli unici termini sono le variabili e le costanti. Sono esempi di formule

$$\forall x \exists y (x \geq y) ; \forall x (x \leq x) ; \forall x (\forall y ((x \leq y) \wedge (y \leq z) \Rightarrow x \leq z)).$$

Linguaggio usato per la teoria dei gruppi. È costituito da "." per rappresentare l'operazione binaria, da "inv" per rappresentare l'operazione che associa ad ogni x il suo inverso e la costante 1 per rappresentare l'elemento neutro. L'unica relazione è l'identità. In generale si preferisce la notazione esponenziale x^{-1} al posto di $\text{inv}(x)$. I termini sono pertanto espressioni del tipo $1 \cdot 1$, $(x \cdot y) \cdot x^{-1}$, $((1 \cdot x)^{-1} \cdot y)^{-1}$. Sono esempi di formule:

$$\forall x (x \cdot y = y \cdot x) ; x \cdot 1 = x ; \forall x (x = 1 \Rightarrow x \cdot x = 1).$$

Ma naturalmente è possibile utilizzare anche linguaggi diversi per trattare lo stesso tipo di strutture matematiche. Ad esempio per la teoria dei gruppi si usa spesso la notazione additiva invece di quella moltiplicativa che abbiamo indicato. In questo caso si utilizza la costante 0 , il nome di operazione binaria $+$, di operazione unaria $-$, ed il solito simbolo di uguaglianza $=$.

Supponiamo di avere un linguaggio del primo ordine $\mathcal{L}$. In accordo con il punto di vista di Hilbert, dobbiamo vedere le dimostrazioni che usualmente si effettuano in matematica come un procedimento meccanico con cui produrre, a partire da un dato sistema di elementi di $\mathcal{L}$ (gli assiomi), nuovi elementi di $\mathcal{L}$ (i teoremi). E-

sistono diverse possibili regole per "produrre teoremi", noi ci soffermiamo sulle seguenti due regole più note in logica.

Il regola del Modus Ponens. Tale regola afferma che se ho dimostrato la formula $\alpha \rightarrow \beta$ ed ho dimostrato α allora posso affermare anche β . In breve:

$$\frac{\alpha, \alpha \rightarrow \beta}{\beta} \quad (\text{Modus Ponens})$$

La regola di Generalizzazione. Tale regola afferma che se ho dimostrato la formula $\alpha(x)$, allora posso affermare anche $\forall x(\alpha)$. In breve

$$\frac{\alpha}{\forall x(\alpha)} \quad (\text{Generalizzazione})$$

La generalizzazione si giustifica col fatto che se ho dimostrato $\alpha(x)$ allora, essendo x una variabile, non ho mai utilizzato nessuna particolare proprietà dell'oggetto denotato da x . In altri termini, durante la dimostrazione x ha sempre denotato un generico elemento del dominio. Pertanto, di fatto, si è dimostrato $\forall x(\alpha)$.

Vi sono poi delle formule di cui ci si può servire durante la dimostrazione perché sono vere sempre, qualunque siano le cose di cui si parla. In altre parole si possono utilizzare delle formule logicamente vere del tipo $\alpha \rightarrow \alpha$ oppure $\alpha \wedge \beta \rightarrow \alpha$ oppure $\alpha(t) \rightarrow \exists x \alpha$. Chiamiamo *assiomi logici* un opportuno insieme A_l di tali formule. Usualmente in A_l si mettono anche formule del tipo

$$p \rightarrow p, \neg(\neg(\alpha)) \rightarrow \alpha, (\forall x \alpha(x)) \rightarrow \alpha(t), \exists x(\alpha) \rightarrow \neg \forall x(\neg(\alpha)).$$

Inoltre si aggiungono assiomi relativi all'uguaglianza

$$U_1 \quad \forall x(x=x) \quad (\text{riflessività})$$

$$U_2 \quad \forall x \forall y((x=y) \rightarrow (y=x)) \quad (\text{simmetria})$$

$$U_3 \quad \forall x \forall y \forall z((x=y) \wedge (y=z)) \rightarrow (x=z) \quad (\text{transitività})$$

$$U_4 \quad x = y \rightarrow z \otimes x = z \otimes y \quad ; \quad (\text{sostituzione})$$

$$x = y \rightarrow x \otimes z = y \otimes z \quad ;$$

$$x = y \rightarrow f(x) = f(y)$$

$$U_5 \quad x = y \rightarrow (r(z,x) \Leftrightarrow r(z,y)) \quad ; \quad (\text{sostituzione})$$

$$x = y \rightarrow (r(x,z) \Leftrightarrow r(y,z)) \quad ;$$

$$x = y \rightarrow (s(x) \Leftrightarrow s(y)).$$

dove U_4 e U_5 devono essere intesi come *schemi di assiomi* cioè devono essere asseriti per tutti i nomi di operazioni binarie $\otimes$, di operazioni unarie f , di relazioni binarie r e di relazioni ad un posto s .

Definizione 5.4. Dato un insieme T di formule, che chiameremo *sistema di assiomi*, chiameremo *dimostrazione di α sotto ipotesi T* ogni successione di formule $\alpha_1, \dots, \alpha_n$ con $\alpha_n = \alpha$ e tale che per ogni formula α_i si verifichi almeno uno dei seguenti casi:

- α_i è un assioma logico
- α_i è una ipotesi, cioè $\alpha_i \in X$
- α_i è stata ottenuta da formule precedenti per modus ponens o per generalizzazione.

Scriveremo $T \vdash \alpha$ per dire che esiste una dimostrazione di α sotto ipotesi T .

Possiamo visualizzare il sistema deduttivo di una logica del primo ordine come una macchina che produce teoremi. Tale macchina, dopo che sono stati inseriti gli assiomi di una teoria T (ad esempio la teoria dei gruppi) comincia a stampare i teoremi uno dopo l'altro. La macchina stampa tali teoremi utilizzando le regole di inferenza e formule che sono in Al , oppure in T oppure che sono state già prodotte.

6. Ma si deve pur parlare di qualche cosa: l'interpretazione

Se ci si attiene strettamente al punto di vista formalista di Hilbert secondo cui la matematica è solo un calcolo di parole, nozioni come essere vero ed essere falso non hanno importanza. L'importante è stabilire regole precise in questa sorta di gioco linguistico, attenersi a tali regole. Tuttavia se ci stacciamo dal riferimento alle "cose di cui si parla" e dalle nozioni di vero e falso non è chiaro come si debbano stabilire tali regole e quale senso abbiano i teoremi trovati.

In definitiva dobbiamo trovare una definizione matematica di interpretazione e di verità. Cominciamo con il problema dell'interpretazione che è alquanto problematico. Ad esempio non è affatto detto che due persone interpretano allo stesso modo una determinata frase. Se si riferisce ad un testo teatrale o un testo musicale è scontato che ogni attore o musicista possa dare una interpretazione personale di tale testo. Ovviamente interpretazio-

ne e verità sono nozioni collegate in quanto una asserzione può essere vera o falsa a seconda del modo come vengono interpretate le parole che la costituiscono. Ad esempio se dico “Maria è più grande di Luisa” la verità o falsità di tale asserzione dipende da chi intendo indicare con i nomi “Maria” e “Luisa” e che cosa intendo per “più grande” (più anziana ?, più alta ?). Anche se scrivo una cosa di carattere matematico come “ $2+2 = 0$ ” scrivo una cosa falsa se interpreto tale asserzione sui numeri naturali, scrivo una cosa vera se la interpreto sugli interi modulo 4. Allora prima di parlare di vero e di falso è necessario fornire una definizione di interpretazione. Ora se $\mathcal{L}$ è un linguaggio del primo ordine, una sua interpretazione si ottiene specificando di quali oggetti si parla e quindi fissando un insieme D . Inoltre, come è ovvio, si deve associare

- ad ogni nome di operazione una operazione in D ,
- ad ogni costante un elemento di D ,
- ad ogni nome di relazione binaria o unaria una relazione binaria o unaria in D .

Definizione 6.1. Una *interpretazione* di un linguaggio del primo ordine $\mathcal{L}$ è costituita da un insieme D , detto *dominio dell'interpretazione* e da una funzione I che associa:

- ad ogni nome di operazione n -aria f una funzione $I(f) : D^n \rightarrow D$
- ad ogni costante c un elemento $I(c)$ di D
- ad ogni nome di relazione n -aria r un sottoinsieme $I(r) \subseteq D^n$.¹⁴

Esempio. Consideriamo un linguaggio $\mathcal{L}$ in cui vi sia solo il nome di una relazione binaria “*ama*”. Allora una interpretazione di tale linguaggio si ottiene fissando un insieme D di persone (ad esempio quelle presenti in una certa stanza) ed una relazione binaria $I(\text{ama})$, cioè un sottoinsieme di $D \times D$. Ad esempio potremmo supporre che D sia l'insieme di persone,

$\{\text{mario, maria, carlo, luigi}\}$

¹⁴ Da notare che la nozione di interpretazione è legata strettamente alla nozione di struttura matematica che abbiamo considerato quando abbiamo parlato dello strutturalismo. Infatti sia dato un linguaggio $\mathcal{L}$ i cui nomi di operazione siano $o_1, \dots, o_s$, i cui nomi di relazioni siano $r_1, \dots, r_t$ e con costanti $c_1, \dots, c_p$. Allora la *struttura matematica associata ad una interpretazione* I è la struttura $(D, I(o_1), \dots, I(o_s), I(r_1), \dots, I(r_t), I(c_1), \dots, I(c_p))$.

e che l'interpretazione sia definita ponendo $I(ama)$ uguale all'insieme

$\{(maria, mario), (mario, carlo), (maria, carlo), (carlo, luigi)\}$.

Naturalmente vi possono essere più interpretazioni dello stesso linguaggio, ad esempio se si cambia il gruppo di persone cui ci si riferisce.

Esempio. Riferiamoci al linguaggio che si usa per le strutture ordinate. Poiché è costituito da un solo simbolo di relazione binaria, $\leq$, una interpretazione di tale linguaggio è costituita da un insieme D e da una relazione binaria $I(\leq)$. Ad esempio possiamo supporre che

- D sia l'insieme dei numeri interi e $I(\leq)$ l'usuale relazione di ordine tra interi,

- D sia l'insieme delle parti di un insieme S e $I(\leq)$ sia la relazione di inclusione.

Si noti che l'unica cosa che si richiede in una interpretazione è che $I(\leq)$ sia una relazione binaria e non necessariamente una relazione d'ordine. Ad esempio otteniamo una interpretazione ponendo D uguale all'insieme dei numeri interi e $I(\leq) = \{(n, m) \mid n = m+1\}$. Naturalmente una tale interpretazione non è un insieme ordinato. Di fatto, almeno che non si impongano opportuni assiomi (cosa che faremo nel seguito) le interpretazioni di questo linguaggio possono essere accettate anche come interpretazioni del linguaggio dell'esempio precedente, e viceversa.

Abbiamo già detto che un termine in cui compaiano delle variabili libere è la descrizione di un algoritmo per calcolare una funzione (se il termine contiene variabili) o un elemento (se il termine non contiene variabili). Ad esempio al termine $3 \cdot (5+7)+1$ corrisponde l'algoritmo:

1. prendi i numeri 5 e 7
2. sommalì
3. moltiplica il risultato per 3
4. aggiungi 1

Chiamiamo interpretazione del termine il numero che si ottiene tramite tale algoritmo. Se nel termine abbiamo la presenza di una o più variabili, allora il corrispondente algoritmo calcola una funzione che chiamiamo interpretazione di tale termine. Naturalmente tali interpretazioni dipendono dalla struttura algebrica di cui si sta parlando, cioè dalla interpretazione fissata. Un modo più rigo-

roso di definire l'interpretazione di un termine è il seguente in cui utilizziamo come variabili i simboli $x_1, x_2, \dots$. Per semplicità supponiamo di utilizzare la notazione prefissa anche per i nomi di operazioni binarie.

Definizione 6.2. Sia (D, I) una interpretazione, e t un termine. Allora per ogni intero n che sia maggiore o uguale agli indici delle variabili che compaiono in t definiamo una funzione n -aria $I(t) : D^n \rightarrow D$ per ricorsione sulla complessità di t al modo seguente:

- a) se t è la costante c allora $I(t)$ è la funzione costantemente uguale a $I(c)$;
- b) se t è la variabile x_i allora $I(t)$ è la proiezione i -esima¹⁵ cioè la funzione definita da

$$I(t)(d_1, \dots, d_n) = d_i$$

- c) se $t = f(\underline{t}_1, \dots, \underline{t}_p)$ allora $I(t)$ è la funzione tale che

$$I(t)(d_1, \dots, d_n) = I(f)(I(\underline{t}_1)(d_1, \dots, d_n), \dots, I(\underline{t}_p)(d_1, \dots, d_n)).$$

Ad esempio sia t il termine $\log(x_1) + \text{sen}(x_2)$, allora

$$I(t)(d_1, d_2) = I(+)(I(\log)(d_1, d_2), I(\text{sen})(d_1, d_2)).$$

D'altra parte

$$I(\log(x_1))(d_1, d_2) = I(\log)(I(x_1)(d_1, d_2)) = I(\log)(d_1)$$

$$I(\text{sen}(x_2))(d_1, d_2) = I(\text{sen})(I(x_2)(d_1, d_2)) = I(\text{sen})(d_2)$$

e quindi

$$I(t)(d_1, d_2) = (I(\log)(d_1) \ I(+)) \ I(\text{sen})(d_2)).$$

Questo significa che il termine $\log(x_1) + \text{sen}(x_2)$ viene interpretato come la funzione che, data la coppia (d_1, d_2) di elementi di D ,

1. applica la funzione $I(\log)$ a d_1 ,
2. applica la funzione $I(\text{sen})$ a d_2 ,
3. compone i risultati ottenuti nei passi 1. e 2. tramite $I(+)$.

Da notare che se in un termine compaiono le variabili x_1 e x_3 allora la sua interpretazione è comunque una funzione di (almeno) tre variabili anche se poi i valori che assume non dipendono effettivamente dai valori di x_2 . Ad esempio il termine $2 \cdot x_1 + \text{sen}(x_3)$ è in-

¹⁵ Prende il nome di *proiezione i-esima* in un prodotto cartesiano $X_1 \times \dots \times X_n$ con $i \leq n$ la funzione Pr_i definita dal porre $Pr_i(x_1, \dots, x_n) = x_i$. In altri termini la proiezione i -esima è la funzione che associa ad ogni vettore la sua componente di posto i . La terminologia geometrica è dovuta al fatto che, ad esempio, in un piano euclideo le due possibili proiezioni corrispondono alle proiezioni sugli assi cartesiani.

interpretato come una funzione di tre variabili in quanto viene considerato equivalente al termine $2 \cdot x_1 + 0 \cdot x_2 + \text{sen}(x_3)$. In tale caso si dice che la variabile x_2 è *muta*. Lo stesso termine può essere interpretato anche come funzione di quattro variabili in quanto può essere considerato equivalente al termine $2 \cdot x_1 + 0 \cdot x_2 + \text{sen}(x_3) + 0 \cdot x_4$. Un termine che si riduca ad una costante può essere interpretato come funzione di n variabili per ogni intero n . Ad esempio se l'interpretazione è il campo dei numeri reali allora la costante 3 può essere interpretata come funzione definita in R in R^2 in R^3 e così via. La cosa non è tanto strana in quanto se, ad esempio, si accetta che la costante 3 sia un termine equivalente al termine $0 \cdot x_1 + 0 \cdot x_2 + 0 \cdot x_3 + 3$ (oppure al termine $0 \cdot x_1 + 0 \cdot x_2 + 3$ oppure al termine $0 \cdot x_1 + 3$) allora tale costante deve poter essere interpretata come funzione di tre variabili (di due variabili e di una variabile, rispettivamente).

7. Cosa è la verità

Allora Pilato gli disse: «Ma dunque, sei tu re?» Gesù rispose: «Tu lo dici; sono re; io sono nato per questo, e per questo sono venuto nel mondo: per testimoniare della verità. Chiunque è dalla parte della verità ascolta la mia voce». Pilato gli disse: «Che cos'è la verità?» E detto questo, uscì di nuovo verso i Giudei ...»
(Giovanni 18:33-38 NRV)

Anche il problema della verità non è tanto semplice tanto che Pilato sembra dubitare perfino che esista la verità. In ogni caso il primo a proporre una definizione di interpretazione e di verità fu il logico matematico Tarski, definizione a cui noi ci atterremo in seguito. Si consideri ora una formula α di un dato linguaggio e proponiamoci di definire in maniera rigorosa che cosa significa l'espressione " α è vera". Ora naturalmente la verità o falsità di una formula dipende dalla interpretazione del linguaggio. Ad esempio la formula $\forall x_1 (\forall x_2 (x_1 \cdot x_2 = x_2 \cdot x_1))$ sarà vera se il dominio D dell'interpretazione è l'insieme degli interi relativi ed il simbolo " $\cdot$ " è interpretato con l'usuale moltiplicazione di numeri. Tale formula sarà invece falsa se invece D è l'insieme delle funzioni di R in R e l'interpretazione di " $\cdot$ " è di essere la composizione di due funzioni (siccome, ad esempio $\log(\text{sen}(x)) \neq \text{sen}(\log(x))$), la

composizione non è una operazione commutativa). Allora è più corretto definire che cosa si debba intendere per

" α è vera rispetto ad una interpretazione I ".

Anche ciò crea qualche difficoltà, infatti se in α vi è una variabile libera x_1 , allora α può essere vera o falsa a seconda dell'elemento rappresentato da x_1 . Ad esempio la formula $\forall x_2(x_2 \cdot x_1 = x_2)$ sarà vera negli interi relativi se x_1 rappresenta l'unità, sarà falsa se x_1 rappresenta il numero 2. Ciò significa che ha senso dire se α è vera o falsa non solo dopo aver fissato una interpretazione I del linguaggio ma anche dopo aver assegnato ad ogni variabile libera di α un particolare elemento nel relativo dominio. Per evitare complicazioni formali nel seguito supponiamo che anche le variabili vincolate siano interpretate da elementi del dominio. Pertanto se le variabili libere o vincolate di α sono tra $x_1, \dots, x_m$, dati $d_1, \dots, d_m$ elementi di D vogliamo dare una definizione precisa del concetto: " α è vera rispetto alla interpretazione I quando le sue eventuali variabili libere sono interpretate con $d_1, \dots, d_m$ ".

Indicheremo in breve con $I \models \alpha [d_1, \dots, d_m]$ una tale asserzione e ne daremo la seguente definizione formale. Per semplicità utilizziamo la notazione prefissa anche per i predicati binari.

Definizione 7.1. Sia I una interpretazione, sia α una formula le cui variabili libere o vincolate siano tra $x_1, \dots, x_m$ e siano $d_1, \dots, d_m$ elementi del dominio D . Allora la relazione $I \models \alpha [d_1, \dots, d_m]$ è definita per induzione sulla complessità di α tramite:

- a) $I \models r(t_1, \dots, t_p) [d_1, \dots, d_m]$ se
 $(I(t_1)(d_1, \dots, d_m), \dots, I(t_p)(d_1, \dots, d_m)) \in I(r)$
- b) $I \models \alpha \wedge \beta [d_1, \dots, d_m]$ se $I \models \alpha [d_1, \dots, d_m]$ e $I \models \beta [d_1, \dots, d_m]$
- c) $I \models \alpha \vee \beta [d_1, \dots, d_m]$ se $I \models \alpha [d_1, \dots, d_m]$ oppure $I \models \beta [d_1, \dots, d_m]$
- d) $I \models \neg \alpha [d_1, \dots, d_m]$ se non è vero che $I \models \alpha [d_1, \dots, d_m]$
- e) $I \models \exists x_i(\alpha) [d_1, \dots, d_m]$ se esiste $d \in D$ tale che
 $I \models \alpha [d_1, \dots, d_{i-1}, d, d_{i+1}, \dots, d_m]$.

Si noti che

$$I \models \forall x_i(\alpha) [d_1, \dots, d_m] \Leftrightarrow I \models \alpha [d_1, \dots, d_i, \dots, d_m] \text{ per ogni } d_i \in D.$$

Infatti, poiché $\forall x_i(\alpha)$ è una abbreviazione di $\neg(\exists x_i(\neg \alpha))$, avremo che

$$\begin{aligned} I \models \forall x_i(\alpha) [d_1, \dots, d_m] \\ \Leftrightarrow \text{non è vero che } I \models \exists x_i(\neg \alpha) [d_1, \dots, d_m] \\ \Leftrightarrow \text{non è vero che esiste } d_i \in D \text{ tale che } I \models \neg \alpha [d_1, \dots, d_i, \dots, d_m] \end{aligned}$$

$\Leftrightarrow$ per ogni $d_i \in D$ $I \models \alpha [d_1, \dots, d_i, \dots, d_m]$.

Definizione 7.2. Se $I \models \alpha [d_1, \dots, d_m]$ diciamo che *la formula α è vera rispetto ad I negli elementi $d_1, \dots, d_m$* . Diciamo che α è vera in I o che I è un modello di α e scriviamo $I \models \alpha$, se risulta $I \models \alpha [d_1, \dots, d_m]$ per ogni $d_1, \dots, d_m$ in D .

In altre parole dire che una formula con eventuali variabili libere è vera in una interpretazione I equivale a dire che la sua chiusura universale è vera. Ad esempio diciamo che $x_1 \cdot x_2 = x_2 \cdot x_1$ è vera in una interpretazione I se $I \models x_1 \cdot x_2 = x_2 \cdot x_1 [d_1, d_2]$ comunque si scelgano d_1 e d_2 nel dominio di interpretazione, cioè se $I \models \forall x_1 \forall x_2 (x_1 \cdot x_2 = x_2 \cdot x_1)$.

8. Teorema di completezza e teoremi limitativi

Siamo ora pronti ad enunciare tre teoremi fondamentali della logica. Il primo parla in positivo sottolineando una cosa (formidabile) che la logica formale riesce a fare. Gli altri due, come vedremo, parlano in negativo mostrando due cose che la logica formale non potrà mai fare.

Abbiamo definito la relazione $\models$ che è di carattere semantico e la relazione $\vdash$ che è di carattere sintattico, cioè relativa ai processi di manipolazione di parole. Il seguente teorema mostra che le due nozioni sono ben collegate.

Teorema 8.1. (Teorema di Completezza¹⁶) Per ogni teoria T ed ogni formula chiusa α risulta:

$$T \models \alpha \Leftrightarrow T \vdash \alpha.$$

Per rendersi conto dell'importanza di tale teorema, osserviamo che la verifica della relazione $T \models \alpha$ richiede l'andare a guardare tutti i possibili modelli di T e nel controllare per ciascun modello se verifica l'asserzione α . Ad esempio se T è la teoria dei gruppi si tratta di verificare che α vale in tutti i possibili gruppi. E' questa una impresa che è impossibile in quanto la classe dei gruppi è tanto grande da condurre, come abbiamo visto, a paradossi. Il teo-

¹⁶ Tale teorema, che è uno dei primi, più importanti ed affascinanti teoremi di logica matematica, è stato dimostrato nel 1929 da Kurt Gödel nella sua tesi di dottorato.

rema di completezza ci dice che verificare $T \vdash \alpha$ è una impresa meno disperata in quanto sembri. Infatti equivale a verificare $T \vdash \alpha$ cioè che la macchina finita che produce passo dopo passo i teoremi di T riesce a produrre anche α . D'altra parte se non valesse il teorema di completezza saremmo nella situazione per cui, ad esempio, esistono asserzioni vere in tutti i gruppi ma che non abbiamo nessuna speranza di poter dimostrare.

Ed ecco gli altri due (famosi) teoremi provati da Gödel nel 1930. Per una loro dimostrazione, come d'altra parte per il teorema di completezza, si rimanda ad un buon testo di logica. Qui ci limitiamo ad enunciarli ed ad esporre le linee generali della dimostrazione del primo teorema per evidenziare come tale dimostrazione dipenda dalla possibilità di autoriferimento di una teoria capace di parlare degli interi. Per semplicità consideriamo una teoria che è formulazione nell'ambito della logica del primo ordine della teoria delle terne di Peano.

Definizione 8.2. Consideriamo un linguaggio con un nome di funzione s ed una costante z_0 . Chiamiamo *teoria del primo ordine per le terne di Peano* il seguente sistema di assiomi:

$$\mathbf{A1} \quad s(x) = s(y) \rightarrow x = y$$

$$\mathbf{A2} \quad \forall x \, z_0 \neq s(x)$$

$$\mathbf{A3} \quad (A(z_0) \wedge \forall x (A(x) \rightarrow A(s(x)))) \rightarrow \forall x A(x).$$

Da notare che A3 è uno “schema di assioma” nel senso che deve essere inteso come rappresentante degli infiniti assiomi che si ottengono fissando in tutti i modi possibili una formula $A(x)$ con una variabile libera x . Si noti che i primi due assiomi coincidono con gli assiomi $P1$ e $P2$ che abbiamo già visto quando abbiamo fornito la nozione di terna di Peano. Invece A3 non coincide con $P3$, infatti asserisce il principio di induzione solo per gli insiemi che sono “descrivibili” tramite una formula $A(x)$ del linguaggio fissato. D'altra parte, poiché esiste solo una quantità numerabile di formule, i sottoinsiemi descrivibili sono in quantità numerabile mentre la classe dei sottoinsiemi di un insieme infinito ha sicuramente potenza maggiore del numerabile. Quindi nella formulazione del primo ordine degli assiomi di Peano il principio di induzione viene richiesto per molti meno insiemi.¹⁷

¹⁷ Naturalmente il fatto che quanto affermato da $P3$ sembra essere più potente da quanto affermato dall'insieme A3 di formule non escludereb-

Teorema 8.3. (Primo Teorema di Gödel) Sia T la teoria del primo ordine delle terne di Peano e supponiamo che sia consistente. Allora esiste una formula ϕ tale che ϕ non può essere né provata né confutata da T .

L'idea per dimostrare tale teorema è ispirata al famoso paradosso che si ottiene considerando l'asserzione

$$\gamma \equiv \text{"io sono una proposizione falsa"}.$$

Allora risulta che

$$\gamma \text{ vera} \Rightarrow \gamma \text{ falsa} ; \gamma \text{ falsa} \Rightarrow \gamma \text{ vera}$$

e pertanto γ non può essere né vera né falsa. Alla base di tale antinomia è il fatto che γ è una proposizione che parla di se stessa, cioè il fenomeno dell' "autoriferimento". Ora una prima "rozza" dimostrazione del primo teorema di Gödel si ottiene partendo dall'asserzione simile

$$\gamma = \text{"io sono una formula che non è un teorema di } T\text{"}.$$

Allora, se si ammette che una tale asserzione sia esprimibile tramite una formula γ del linguaggio $\mathcal{L}$,

- se $T \vdash \gamma$ allora γ non è un teorema di T ;

- se $T \vdash \neg\gamma$ allora γ è un teorema di T

Pertanto, come volevamo dimostrare, né γ né la sua negata $\neg\gamma$ possono essere teoremi di T .

be in linea di principio che si possa dimostrare l'equivalenza delle due formulazioni. In altre parole non escluderebbe che il principio di induzione per la classe dei sottoinsiemi descrivibili possa implicare il principio di induzione per tutti i sottoinsiemi. Per potere convincersi che non sussiste l'equivalenza dovremmo infatti mostrare che esiste un modello di $A1, A2, A3$ che non verifica $P1, P2, P3$. Il materiale per trovare un tale modello lo possiamo trovare nel campo di razionali non-standard che abbiamo ottenuto come ultra-potenza di $\mathbb{Q}$. Infatti chiamiamo *numero naturale non-standard* un elemento $[(r_n)_{n \in \mathbb{N}}]$ di $\mathbb{Q}^*$ con $(r_n)_{n \in \mathbb{N}}$ successione di numeri naturali. Detto N_0^* l'insieme di tali numeri, definiamo la funzione $s : N_0^* \rightarrow N_0^*$ ponendo $s([(r_n)_{n \in \mathbb{N}}]) = [(r_n + 1)_{n \in \mathbb{N}}]$. Allora $(N_0^*, s, 0)$ verifica $A1, A2, A3$ ma non verifica $P3$. Infatti per un teorema fondamentale della teoria delle ultrapotenze, $(N_0^*, s, 0)$ verifica tutte le proprietà dell'insieme dei naturali che siano esprimibili al primo ordine. In particolare $(N_0^*, s, 0)$ verifica $A1, A2, A3$. D'altra parte l'insieme degli interi che non sono infiniti, pur contenendo 0 e pur essendo chiuso per successore, non coincide con N_0^* . Quindi non verifica $P3$.

Per potere formalizzare quanto sopra detto, è necessario che sia possibile il fenomeno dell' autoriferimento, abbiamo cioè bisogno di far vedere come T possa "parlare di se stesso". In particolare, perché quel termine "io" abbia senso deve essere dato un nome all'interno del linguaggio $\mathcal{L}$ a ciascuna delle formule di $\mathcal{L}$. Inoltre deve essere definita una formula in $\mathcal{L}$ che significhi "essere teorema". Per fare ciò cominciamo con il mettere in rilievo che:

ogni numero naturale ha un "nome" corrispondente in $\mathcal{L}$.

Infatti nel linguaggio delle terne di Peano, che contiene il simbolo s di successivo, possiamo dare un nome ad ogni numero denotando con il termine chiuso $s(0)$ il numero uno, con $s(s(0))$ il numero due e così via. Inoltre

è possibile codificare le formule di $\mathcal{L}$.

Cioè è possibile associare ad ogni formula ϕ un numero intero detto *numero di codice* di ϕ . Infatti l'insieme delle formule di una logica del primo ordine è un insieme di parole ed abbiamo già affrontato il problema della codifica di un insieme di parole. Da ciò segue che:

Ad ogni formula ϕ di $\mathcal{L}$ può essere assegnato un "nome" $c(\phi)$ in $\mathcal{L}$ considerando il termine chiuso $c(\phi)$ che rappresenta il numero di codice di ϕ .

Similmente si vede che

Si può dare un numero di codice ad ogni dimostrazione π in T .

La cosa non è difficile poiché una dimostrazione può essere vista come una sequenza di formule $\alpha_1, \alpha_2, \dots, \alpha_n$ e tale sequenza è una parola nell'alfabeto che si ottiene aggiungendo all'alfabeto A il simbolo $"^{18}$. Si può pertanto procedere allo stesso modo di quanto si è fatto per la codifica delle formule.

¹⁸ Per motivi tecnici che qui sarebbe troppo lungo spiegare la codifica delle dimostrazioni deve essere un processo effettivo e questa effettività richiede che l'insieme degli assiomi della teoria sia decidibile, cioè che

Ad ogni dimostrazione π può essere assegnato un "nome" $c(\pi)$.

Come nel caso delle formule basta considerare il termine chiuso che rappresenta il numero di codice di π .

Detto questo si dimostra (ma noi non lo dimostriamo) che in $\mathcal{L}$ esiste una formula $Pr(x,y)$ il cui significato è che x è (il numero di codice di) una dimostrazione di y (della formula codificata da y). Più precisamente Pr verifica la seguente proprietà

$$"T \vdash \phi \text{ se e solo se esiste un termine chiuso } t \text{ tale che } T \vdash Pr(t, c(\phi))."$$

Infine si prova l'esistenza di una formula γ tale che

$$T \vdash \gamma \leftrightarrow (\neg \exists x Pr(x, c(\gamma))).$$

La formula γ asserisce proprio quello che volevamo, cioè che "io sono una formula che non è un teorema di T ". Supponiamo ora che γ sia dimostrabile, allora sarebbe dimostrabile anche $\neg \exists x Pr(x, c(\gamma))$ e quindi non potrebbe esistere una dimostrazione di γ in T . Supponiamo invece che $\neg \gamma$ sia dimostrabile, allora sarà dimostrabile in T anche $\exists x Pr(x, c(\gamma))$. Ciò comporta che esiste un termine chiuso t per cui $Pr(t, c(\gamma))$ e pertanto che esiste una dimostrazione di γ . Ciò è in contrasto con l'ipotesi di consistenza per T .

Corollario 8.4. Se T è consistente, allora esiste una asserzione dell'aritmetica che pur essendo vera non può essere dimostrata. In altre parole T non è abbastanza potente da permettere di provare tutte le proposizioni vere dell'aritmetica.

Dim. Detta ϕ la formula indecidibile, se si ammette il modello naturale dell'aritmetica, allora in tale modello sarà vera ϕ oppure $\neg \phi$. Nel primo caso ϕ è una proposizione vera che non può essere dimostrata, nel secondo caso la stessa cosa si può dire per $\neg \phi$.

Teorema 8.5. (Secondo Teorema di Gödel). Se T è consistente, allora la formula $\neg \exists x Pr(x, \gamma \wedge \neg \gamma)$ che asserisce la consistenza di T non si può provare in T .

esista un procedimento effettivo per capire se una formula è un assioma oppure no.

Da notare che anche per tale teorema la possibilità di autoriferimento è essenziale.

Quanto ora dimostrato non vale solo per la teoria del primo ordine delle terne di Peano ma per ogni teoria che sia abbastanza potente da permettere di definire i numeri naturali e quindi per ogni teoria che pretenda di fondare la matematica. Ad esempio vale anche per tutte le teorie che vogliano assiomatizzare la nozione di insieme. In particolare, il secondo teorema di Gödel mostra che qualunque sia tentativo di fornire un adeguato sistema di assiomi T per la matematica, la consistenza di T non può essere provata all'interno della stessa teoria T . In altri termini per provare la consistenza di T dobbiamo necessariamente utilizzare strumenti più potenti di T stesso.

Non è quindi possibile, come sperava Hilbert, provare la consistenza di teorie "forti" che coinvolgono l'infinito attuale tramite metodi finitisti.

APPENDICE

NOZIONI BASE E VARIE

1. Coppie, prodotti cartesiani e relazioni

Questo libro è rivolto a persone che abbiano già una conoscenza elementare della matematica. Tuttavia, per permetterne una lettura ad una platea la più ampia possibile, in questa appendice ricordo alcune semplici nozioni utilizzate nel libro. Nel seguito indicherò:

- con $\emptyset$ l'insieme vuoto,
- con $\{d_1, \dots, d_n\}$ l'insieme i cui elementi sono $d_1, \dots, d_n$.
- con $\{x : x \text{ verifica la proprietà } P\}$ l'insieme i cui elementi sono tutti e soli quelli verificanti la proprietà P .

Il primo passo per la definizione dei concetti fondamentali della teoria degli insiemi è quello di definire la nozione di coppia ordinata e di prodotto cartesiano.

Definizione 1.1. Dati due elementi x ed y chiamiamo *coppia di primo elemento x e secondo elemento y* l'insieme $\{\{x\}, \{x, y\}\}$ che indichiamo con (x, y) . Dati due insiemi X ed Y chiamiamo *prodotto cartesiano* l'insieme $X \times Y = \{(x, y) : x \in X, y \in Y\}$ delle coppie costituite da un elemento di X ed un elemento di Y .

Da notare che gli insiemi $\{\{x\}, \{x, y\}\}$, $\{\{x, y\}, \{x\}\}$, $\{\{y, x\}, \{x\}\}$ e $\{\{x\}, \{x, y\}\}$ coincidono tra loro e rappresentano tutti la stessa coppia (x, y) . Inoltre la coppia (x, x) è rappresentata dall'insieme $\{\{x\}\}$.

Definizione 1.2. Chiamiamo relazione *binaria* tra un insieme X ed un insieme Y ogni sottoinsieme $\mathcal{R}$ di $X \times Y$. L'*inversa* di $\mathcal{R}$ è la relazione $\mathcal{R}^{-1} = \{(x, y) : (y, x) \in \mathcal{R}\}$. Date due relazioni $\mathcal{R}_1 \subseteq X \times Y$ e $\mathcal{R}_2 \subseteq Y \times Z$, la loro *composizione* è la relazione $\mathcal{R}_1 \circ \mathcal{R}_2$ definita ponendo

$$\mathcal{R}_1 \circ \mathcal{R}_2 = \{(x, z) \in X \times Z : \text{esiste } y \in Y, (x, y) \in \mathcal{R}_1, (y, z) \in \mathcal{R}_2\}.$$

Definizione 1.3. Data una relazione $\mathcal{R}$ il suo *dominio* viene definito ponendo

$$\text{Dom}(\mathcal{R}) = \{x \in X : \text{esiste } y \in Y \text{ per cui } (x, y) \in \mathcal{R}\},$$

il suo *codominio* ponendo

$$\text{Cod}(\mathcal{R}) = \{y \in Y : x \in X \text{ per cui } (x,y) \in \mathcal{R}\}.$$

Usualmente per le relazioni binarie si utilizza la *notazione infissa* che consiste nello scrivere $x\mathcal{R}y$ per indicare che $(x,y) \in \mathcal{R}$.

Definizione 1.4. Una relazione binaria $f \subseteq X \times Y$ è chiamata *funzione di X in Y* se è univoca, cioè se, per ogni $x \in X$ esiste al più un elemento y tale che $(x,y) \in f$. Se per ogni $x \in X$ esiste uno ed un solo elemento y tale che $(x,y) \in f$, cioè se $\text{Dom}(f) = X$ allora diciamo che f è *totale*, altrimenti che è *parziale*.

In analisi matematica $\text{Dom}(f)$ viene chiamato anche *campo di esistenza* di f . Nel seguito denoteremo una funzione con una lettera minuscola, ad esempio la lettera f e scriveremo $f: X \rightarrow Y$ per indicare che f è una funzione di X in Y . Per ogni $x \in X$ indichiamo con $f(x)$ l'unico elemento y tale che $(x,y) \in f$. In alcuni testi l'insieme X viene detto *dominio*, l'insieme Y viene detto *codominio* di f . In altri testi per dominio e codominio si intendono i due insiemi $\text{Dom}(f)$ e $\text{Cod}(f)$.

Definizione 1.5. Una funzione $f: X \rightarrow Y$ è *iniettiva* se

$$f(x) = f(y) \Rightarrow x = y,$$

è *suriettiva* se per ogni $y \in Y$ esiste almeno un $x \in X$ tale che $f(x) = y$ cioè se $\text{Cod}(f) = Y$, è *biettiva* se è sia suriettiva che iniettiva.

2. Definizione (brutta) di n -pla

Anche se in tutti i libri viene proposta la definizione di coppia che abbiamo dato nel paragrafo precedente, ci si pone il problema di perché si sia optato per una definizione così strana, arbitraria e comunque poco intuitiva. Ancora meno intuitiva è la nozione di n -pla che viene fatta per induzione su $n \geq 2$ al modo seguente.

Definizione 2.1. Dato $n \geq 2$ e gli insiemi $X_1, \dots, X_n$ definiamo il prodotto cartesiano di tali n insiemi per ricorsione su n tramite l'equazione

$$X_1 \times \dots \times X_n = (X_1 \times \dots \times X_{n-1}) \times X_n$$

e chiamando n -pla un elemento di tale prodotto cartesiano.

Pertanto una terna $(x,y,z) = ((x,y),z)$ viene a coincidere con l'insieme

$$\{\{\{x\},\{x,y\}\},\{\{\{x\},\{x,y\}\},z\}\}.$$

In particolare la terna (x,x,x) coincide con l'insieme $\{\{\{x\}\},\{\{\{x\}\},x\}\}$. Lascio a chi legge il divertimento di dire che cosa è una quadrupla (x,x,x,x) . Solo i matematici sono capaci di rappresentare una cosa tanto semplice in modo tanto tortuoso ! La questione è che i matematici hanno introdotto la teoria degli insiemi per il desiderio di trovare uno strumento unico per costruire tutta la matematica. Quindi sono costretti a rappresentare quella che è intuitivamente una coppia come $(2,5)$ usando solo strumenti insiemistici e quindi solo le parentesi $\{, \}$ ed i simboli 2,5. La rappresentazione deve essere tale da rendere possibile la “estrazione” l'informazione di quale si intenda come primo elemento e quale come il secondo. La Definizione 1.1 permette appunto di ottenere questo. La nozione di n -pla per $n > 2$ può essere invece semplificata ricorrendo a quella di funzione. Ad esempio una terna potrebbe essere definita come una funzione definita in $\{1,2,3\}$ dopo avere indicato con 1, 2, 3 gli insiemi $\{\emptyset\}$, $\{\{\emptyset\}\}$, $\{\{\{\emptyset\}\}\}$. D'altra parte quando pensiamo ad una n -pla $(x_1, \dots, x_n)$ pensiamo appunto ad una corrispondenza che associa ad ogni “posto” in $\{1, \dots, n\}$ un elemento. Se $X_1, \dots, X_n$ sono insiemi allora il prodotto cartesiano verrebbe definito ponendo

$$X_1 \times \dots \times X_n = \{f : \{1, \dots, n\} \rightarrow X_1 \cup \dots \cup X_n \text{ tali che } f(i) \in X_i\}.$$

Una tale definizione si presta bene ad essere generalizzata al caso infinito al modo seguente.

Definizione 2.2. Sia $(S_i)_{i \in I}$ una famiglia di insiemi. L'insieme delle applicazioni $f : I \rightarrow \bigcup_{i \in I} S_i$ tali che $f(i) \in S_i$, per ogni $i \in I$, viene detto *prodotto cartesiano* della famiglia $(S_i)_{i \in I}$ e lo si indica con $\times_{i \in I} S_i$.

Fortunatamente i matematici dopo avere imparato la definizione insiemistica di n -pla la dimenticano rapidamente per tornare a basarsi, nei propri ragionamenti, sulla nozione intuitiva che tutti abbiamo.

3. Relazioni di equivalenza e quozienti

Esistono proprietà particolarmente importanti per le relazioni binarie su di un dato insieme. Ne elenchiamo alcune:

Definizione 3.1. Dato un insieme S ed una relazione binaria $\mathcal{R}$ in S diremo che

- $\mathcal{R}$ è *riflessiva* se $x\mathcal{R}x$ per ogni $x \in S$
- $\mathcal{R}$ è *transitiva* se per ogni $x, y, z \in S$
 $x\mathcal{R}y$ e $y\mathcal{R}z \Rightarrow x\mathcal{R}z$
- $\mathcal{R}$ è *simmetrica* se per ogni $x, y \in S$
 $x\mathcal{R}y \Rightarrow y\mathcal{R}x$
- $\mathcal{R}$ è *asimmetrica* se per ogni $x, y \in S$
 $x\mathcal{R}y$ e $y\mathcal{R}x \Rightarrow x = y$
- $\mathcal{R}$ è *totale* o *lineare* se per ogni $x, y \in S$
 $x\mathcal{R}y$ oppure $y\mathcal{R}x$.

Definizione 3.2. Una relazione binaria $\mathcal{R}$ in un insieme S è una *relazione di equivalenza* se è riflessiva, transitiva e simmetrica.

In generale indicheremo con $\equiv$ una relazione di equivalenza.

Esercizio. Dato un insieme non vuoto S , dimostrare che l'identità, cioè la relazione

$$D(S) = \{(x, y) : x \text{ ed } y \text{ sono lo stesso elemento di } S\}$$

è una relazione di equivalenza (la più piccola su S).

Esercizio. Dimostrare che la relazione totale, cioè la relazione $\mathcal{R} = S \times S$ secondo cui tutti gli elementi sono equivalenti tra loro è una relazione di equivalenza (la più grande su S).

Le relazioni di equivalenza sono alla base delle *definizioni per astrazione* in cui si considerano come un unico oggetto oggetti che, pur essendo diversi, differiscono per aspetti che si considerano non essenziali. Ad esempio consideriamo il seguente enunciato di un problema:

calcolare l'area di un triangolo di lati 3, 4, 5.

Spesso lo stesso problema viene enunciato al modo seguente:

calcolare l'area del triangolo di lati 3, 4, 5.

Questa seconda formulazione del problema è corretta ? Il fatto che esistono infiniti triangoli con lati 3,4,5 sembra mostrare che non lo sia. Tuttavia l'uso del singolare significa che si è deciso di considerare un nuovo oggetto "astratto" il triangolo di lati 3, 4, 5 che in un certo senso rappresenta tutti i possibili triangoli di lati 3, 4, 5. Ciò è possibile poiché, per quanto riguarda il problema da affrontare, non è interessante sapere la posizione di un triangolo sul piano ma solo le sue dimensioni. Alla base di tale processo è la relazione di equivalenza per cui due triangoli sono *uguali* se hanno lati uguali. Per le strutture algebriche si procede in modo analogo. E possibile dire un gruppo di ordine 5 ma è possibile dire anche il gruppo di ordine 5. Infatti tutti i gruppi di ordine 5 sono isomorfi tra loro e l'isomorfismo è una relazione di equivalenza.

Un modo per formalizzare un tale modo di procedere è identificare un oggetto astratto definito in questo modo con l'insieme degli oggetti *concreti* da esso rappresentato.

Definizione 3.3. Data una relazione di equivalenza $(S, \equiv)$ ed $x \in S$, la *classe completa di equivalenza determinata da x* è definita ponendo

$$[x] = \{x' \in S : x' \equiv x\}.$$

Il *quoziente di S modulo $\equiv$* è l'insieme $S/\equiv$ delle classi complete di equivalenza, cioè

$$S/\equiv = \{[x] : x \in S\}.$$

In altre parole se parto da un universo di oggetti S ed introduco una relazione di equivalenza $\equiv$ tra oggetti di tale insieme, allora vengo a creare per astrazione un nuovo insieme di oggetti $S/\equiv$. Tornando all'esempio dei triangoli, l'espressione "*calcolare l'area del triangolo di lati 3, 4, 5*" diventa corretta se col termine "triangolo di lati 3, 4, 5" si intende un unico oggetto: la classe completa di equivalenza costituita da tutti i triangoli i cui lati hanno tali lunghezze.

Proposizione 3.4. Siano S ed S' due insiemi ed $f : S \rightarrow S'$ una funzione. Allora la relazione $\equiv$ definita ponendo

$$x \equiv y \Leftrightarrow f(x) = f(y)$$

è una relazione di equivalenza che viene detta il *nucleo* di f . Inoltre ogni relazione di equivalenza si ottiene in questo modo, cioè ogni relazione di equivalenza è il nucleo di qualche funzione.

Dim. La prima parte della dimostrazione si lascia per esercizio. Sia $\equiv$ una qualunque relazione di equivalenza in un insieme S e sia $S' = S/\equiv$. Allora la funzione $f: S \rightarrow S'$ ottenuta ponendo $f(x) = [x]$ è tale che

$$f(x) = f(y) \Leftrightarrow [x] = [y] \Leftrightarrow x \equiv y. \quad \square$$

Nell'esempio dei triangoli la relazione di eguaglianza è determinata dalla funzione che associa ad ogni triangolo la terna costituita dalla lunghezza dei suoi lati.

Esempio: Sia S l'insieme i cui elementi sono mucchietti di monete e sia f la funzione che associa ad ogni mucchietto x la somma totale rappresentata da x . Allora due mucchietti sono da considerare equivalenti se corrispondono allo stesso valore.

4. Relazioni d'ordine e reticoli

Una classe importante di relazioni binarie in un insieme sono le relazioni d'ordine.

Definizione 4.1. Diciamo che una relazione binaria $\mathcal{R}$ in un insieme S è una *relazione di pre-ordine* se è riflessiva e transitiva. Diciamo che $\mathcal{R}$ è una *relazione d'ordine* se è riflessiva, transitiva ed asimmetrica.

Le relazioni di pre-ordine e di ordine usualmente vengono denotate con il simbolo $\leq$. Se S è un insieme ed $\mathcal{R} \subseteq S \times S$ una relazione binaria su S allora a volte si usa dire che la coppia $(S, \mathcal{R})$ è una *struttura relazionale*.

Definizione 4.2. Date due strutture relazionali $(S_1, \mathcal{R}_1)$ e $(S_2, \mathcal{R}_2)$ chiamiamo *isomorfismo* di $(S_1, \mathcal{R}_1)$ in $(S_2, \mathcal{R}_2)$ ogni funzione biettiva $f: S_1 \rightarrow S_2$ tale che

$$(x, y) \in \mathcal{R}_1 \Leftrightarrow (f(x), f(y)) \in \mathcal{R}_2.$$

Definizioni 4.3. Data una struttura di preordine $(S, \leq)$ ed un sottoinsieme X di S prende il nome di *maggiorante* (*minorante*) di X un elemento $m \in S$ tale che $x \leq m$ ($m \leq x$) per ogni $x \in X$. Prende il nome di *massimo* (*minimo*) *elemento* di X un elemento $m \in X$ tale che $m \geq x$ ($x \geq m$) per ogni $x \in X$. Si chiama *estremo superiore* un elemento $\sup(X)$ che sia il minimo dell'insieme dei maggioranti di X . Si chiama *estremo inferiore* un elemento $\inf(X)$ che sia il massimo dell'insieme dei minoranti di X .

Non è detto che gli estremi superiori o inferiori esistano sempre. Ad esempio l'insieme dei numeri primi non ammette estremo superiore.¹ Se m è il minimo di X è anche l'estremo inferiore ma il viceversa non vale. La stessa cosa si può dire per il massimo. Ad esempio se X è l'intervallo aperto $(0,1)$ allora 0 è l'estremo inferiore ma non è il minimo in quanto non appartiene ad X . 1 è l'estremo superiore ma non è il massimo.

Definizione 4.4. Si chiama *reticolo* un insieme ordinato $(S, \leq)$ tale che per ogni coppia x ed y di elementi di S esistono $\sup\{x, y\}$ e $\inf\{x, y\}$. Si dice che $(S, \leq)$ è un *reticolo completo* se per ogni insieme X di elementi di S esistono $\sup(X)$ e $\inf(X)$.

Dato un reticolo completo esiste l'estremo superiore della famiglia di tutti gli elementi di S . Tale estremo superiore è ovviamente il massimo di S ed usualmente viene denotato con 1 . Similmente esiste l'estremo inferiore della famiglia di tutti gli elementi di S . Tale estremo inferiore è il minimo di S e viene usualmente denotato con 0 .

In un reticolo possiamo definire due operazioni $\wedge$ e $\vee$ ponendo

$$x \wedge y = \inf\{x, y\} \text{ e } x \vee y = \sup\{x, y\}.$$

In tale modo si ottiene una struttura algebrica $(S, \vee, \wedge)$ verificante le seguenti proprietà.

- (i) $x \vee y = y \vee x$; $x \wedge y = y \wedge x$ (commutativa)
- (ii) $x \vee (y \vee z) = (x \vee y) \vee z$; $x \wedge (y \wedge z) = (x \wedge y) \wedge z$ (associativa)
- (iii) $x \vee x = x$; $x \wedge x = x$ (idempotenza)
- (iv) $x \leq y \Leftrightarrow x \wedge y = x$.

¹ Ricordiamo che gli antichi greci provarono che per ogni numero primo p esiste un numero primo q maggiore di p (in termini attuali diremmo che l'insieme dei numeri primi è infinito).

Viceversa è possibile provare che se si considera una struttura algebrica $(S, \vee, \wedge)$ verificante (i), ii) e iii), allora posto per definizione $x \leq y$ se $x \wedge y = x$, la struttura $(S, \leq)$ è un reticolo. In altre parole è possibile introdurre i reticoli sia come strutture d'ordine che come strutture algebriche.

Definizione 4.5 (Definizione algebrica). Chiamiamo *reticolo* una struttura algebrica $(S, \vee, \wedge)$ che verifica i), ii) e iii).

5. Relazioni di buon ordine

Tra le relazioni d'ordine sono di particolare importanza le relazioni di “buon ordine”.

Definizione 5.1. Una relazione di ordine $(S, \leq)$ è detta di *buon ordinamento* se ogni sottoinsieme non vuoto di S ammette minimo.

Nel terzo capitolo abbiamo visto che l'ordinamento che si definisce in una terna di Peano (cioè nell'insieme dei numeri naturali) è un esempio di insieme infinito con un buon ordinamento.

$\{0, 1, 2, 3, 4, \dots\}$

Poiché è evidente che ogni sottoinsieme di un insieme con un buon ordinamento è un insieme con un buon ordinamento, ogni insieme finito di numeri naturali è un insieme ben ordinato. In particolare sono esempi di insiemi ben ordinati gli insiemi finiti

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

Naturalmente ci stiamo riferendo all'ordinamento usuale in tali insiemi. Se si considerano le lettere $\{a, b, c, d\}$ con l'usuale ordinamento delle lettere dell'alfabeto abbiamo un esempio di insieme finito con un buon ordinamento.

Definizione 5.2. Dato un insieme ben ordinato $(S, \leq)$ chiamiamo *successivo* di un elemento x un elemento $s(x)$ tale che $x < s(x)$ e non esiste z per cui $x < z < s(x)$. Un elemento, diverso da 0, che non è successivo di nessun elemento viene chiamato *elemento limite*.

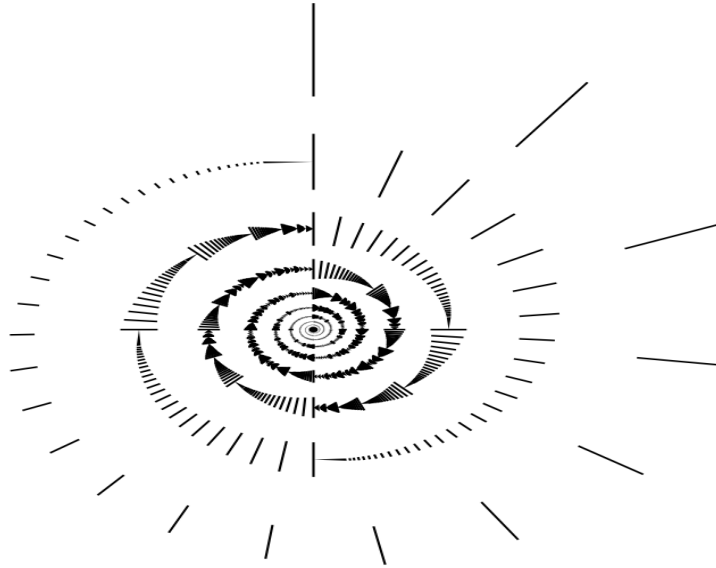
Negli ordinali finiti ed in una terna di Peano non esistono elementi limite.

Esercizio. Dimostrare che se

$S = \{1-1/n : n \in \mathbb{N}\} \cup \{2-1/n : n \in \mathbb{N}\} \cup \{3-1/n : n \in \mathbb{N}\}$
 e $\leq$ è l'ordine usuale tra numeri reali, allora $(S, \leq)$ è un insieme ben ordinato in cui i numeri 1 e 2 sono elementi limite.

Problema. Provare che una struttura del tipo $(\mathcal{P}(S), \subseteq)$ non è di buon ordine. Provare inoltre che l'anello $\mathbb{Z}$ degli interi, il campo $\mathbb{Q}$ dei razionali o il campo dei reali non sono di buon ordine.

La seguente immagine, presa da Wikipedia, rappresenta un ordinale in cui esistono infiniti elementi limite. Si lascia al lettore il compito di interpretare la figura e di individuare gli elementi limite:



Un insieme con un buon ordine assomiglia per molti aspetti ad una terna di Peano.

Proposizione 5.3. Ogni relazione di buon ordine $(S, \leq)$ è una relazione d'ordine totale con elemento minimo. Inoltre ogni elemento x in S che non sia il massimo ammette un *successivo*.

Dim. Per provare che l'ordine è totale consideriamo due elementi x ed y . Allora poiché l'insieme $X = \{x, y\}$ per ipotesi è dotato di minimo e questo minimo o è x oppure y . Nel primo caso risulterà

$x \leq y$, nel secondo caso $y \leq x$. Inoltre poiché ogni sottoinsieme di S ammette minimo in particolare anche S ammette minimo. Infine, dato $x \in S$ che non sia il massimo poniamo $s(x)$ uguale al minimo dell'insieme (non vuoto) $\{z \in S : x < z\}$. Allora è evidente che tra x ed $s(x)$ non possono esserci elementi. $\square$

Un esempio più interessante di ordinamento che non è un buon ordinamento si ottiene considerando l'usuale ordinamento nell'insieme $Q^+ = \{x \in Q : x \geq 0\}$ dei numeri razionali non negativi. Infatti il sottoinsieme $\{1/n : n \in N\}$ di Q^+ non ammette minimo (pur ammettendo 0 come estremo inferiore). Come vedremo, niente esclude che in Z ed in Q possano essere definiti dei buoni ordinamenti diversi da quello naturale.

Ricordiamo che un isomorfismo tra due insiemi ordinati $(S, \leq)$ ed $(S', \leq')$ è una funzione biettiva tale che

$$x \leq y \Leftrightarrow f(x) \leq' f(y).$$

Negli insiemi con una relazione di buon ordine è possibile nelle dimostrazioni l'utilizzazione di una estensione del principio di induzione che prende il nome di *principio di induzione transfinita*.

Teorema 5.4. (Principio di induzione transfinita). Consideriamo un insieme ben ordinato $(S, \leq)$ ed una proprietà P definita in S . Allora se è verificata l'implicazione

i) P vale per ogni $x < y \Rightarrow P$ vale per y
possiamo concludere che

ii) P vale per ogni $x \in S$.

Dim. Assumiamo che valga i), allora per provare ii) proviamo che l'insieme $X = \{x \in S : P \text{ è falsa in } x\}$ è vuoto. Infatti se non fosse vuoto ammetterebbe un minimo x_0 . Per definizione di minimo ciò significa che $x_0 \in X$ e $x \notin X$ per ogni $x < x_0$. In altre parole P sarebbe falsa in x_0 mentre sarebbe vera per ogni $x < x_0$. Ciò è in contrasto con i). $\square$

Poiché nelle terne di Peano è definito un buon ordine, in tali strutture è possibile effettuare dimostrazioni con tale principio che a volte può risultare più comodo. Ad esempio, esso viene utilizzato per la dimostrazione del seguente famoso teorema.

Teorema 5.5. (Teorema fondamentale dell'aritmetica) Ogni numero naturale n diverso da zero può essere scomposto in un unico modo nel prodotto di un numero finito di primi.

Dim. Trascurando l'unicità, vogliamo dimostrare solo che n si può scrivere come prodotto di un numero finito di primi. Potremmo tentare di utilizzare il principio di induzione, ma in tale caso dovremmo provare che se n è scomponibile allora $n+1$ è ancora scomponibile: cosa questa che non sembra facile perché sembra difficile ricavare da una scomposizione di n una scomposizione di $n+1$. Proviamo invece ad utilizzare il principio di induzione transfinita. Supponiamo pertanto che tutti gli interi strettamente minori di n ammettono una scomposizione e proponiamoci di dimostrare che anche n ammette una scomposizione. Ora se n è primo allora banalmente ammette una scomposizione in fattori primi. Se n non è primo sarà uguale al prodotto di due interi a e b strettamente minori di n i quali pertanto, per ipotesi di induzione, sono scomponibili. Pertanto n , essendo prodotto di due numeri scomponibili, è ancora scomponibile.

Naturalmente si pone il problema della esistenza di buoni ordinamenti. Come abbiamo visto le terne di Peano sono esempi di insiemi con un buon ordinamento. Altri esempi si ottengono tramite la proposizione seguente.

Proposizione 5.6. Sia $(S, \leq)$ un buon ordinamento e sia $(S', \leq')$ una struttura binaria tale che esista una immersione $f: S' \rightarrow S$ di $(S', \leq')$ in $(S, \leq)$. Allora $(S', \leq')$ è un buon ordinamento.

Dim. Consideriamo un sottoinsieme non vuoto X' di S' . Allora $f(X')$ è un sottoinsieme non vuoto di S e quindi ammette un minimo m' . Detto $m \in X$ tale che $f(m) = m'$, risulta che m è il minimo di X . Infatti se $x' \in X$, allora essendo $f(m) = m' \leq f(x')$, ed essendo f una immersione, siamo sicuri che $m \leq m'$.

Corollario 5.7. Ogni struttura isomorfa ad in insieme con un buon ordine è un insieme con un buon ordine.² Ogni sottoinsieme di un insieme ben ordinato è un insieme ben ordinato.

² Tale asserzione è un caso particolare del fatto che se due strutture sono isomorfe allora verificano le stesse proprietà. Ne segue che se una strut-

Dim. Osserviamo solo che se X è un sottoinsieme di un insieme ben ordinato S allora l'immersione identica è una immersione di X in S .

Proposizione 5.8. In ogni insieme numerabile è possibile definire un buon ordinamento isomorfo a quello dei numeri naturali. In particolare in Z ed in Q è definibile un buon ordinamento.³

Dim. Supponiamo che S sia un insieme numerabile, allora esiste una funzione biettiva $f: N \rightarrow S$ di N in S . Definiamo in S la relazione $\leq'$ ponendo

$$f(n) \leq' f(m) \Leftrightarrow n \leq m.$$

Allora è evidente che $(S, \leq')$ è una struttura relazionale isomorfa ad $(N, \leq)$ e che f è un isomorfismo tra tali strutture. Dalla Proposizione 5.7 segue che $(S, \leq')$ è un buon ordinamento. $\square$

In altre parole, una volta fissata una strategia di enumerazione degli elementi di S diciamo che un elemento $x \in S$ è minore di un elemento $y \in S$ se in tale enumerazione in un certo senso x “viene prima” di y . Ad esempio, supponiamo di aver deciso di enumerare gli elementi di Z iniziando da 0 e poi alternando i negativi con i positivi

$$0, -1, +1, -2, +2, \dots$$

Avremo allora che si ottiene un buon ordinamento in Z al modo seguente:

$$0 <' -1 <' +1 <' -2 <' +2 \dots$$

E' interessante osservare che in Z è possibile definire anche un buon ordinamento che non è isomorfo a quello di N . Infatti basta ordinare gli elementi di Z ponendo prima 0, poi tutti i negativi e poi tutti i positivi al modo seguente:

tura è un modello di un dato sistema di assiomi anche la seconda è un modello dello stesso sistema di assiomi. Una tale proprietà per due strutture in logica matematica viene espressa dall'asserzione “se due strutture sono isomorfe allora sono logicamente equivalenti”. Ad esempio se una struttura algebrica è isomorfa ad un gruppo G allora è ancora un gruppo, se G è commutativo allora tale struttura è un gruppo commutativo.

³ Più complesso il discorso per l'insieme dei numeri reali R o, più in generale, per un qualsiasi insieme. Infatti in tale caso la dimostrazione di esistenza di un buon ordine è possibile solo se si utilizza l'assioma della scelta.

$$0 < -1 <^* -2 <^* -3 <^* \dots <^* 1 <^* 2 <^* 3 \dots$$

Una tale relazione d'ordine è quella dell'ordinale $\omega + \omega$ (per la nozione di ordinale si veda il Capitolo 4). Un discorso simile può essere fatto anche per l'insieme Q^+ dei razionali positivi che abbiamo già dimostrato essere numerabile. Basta vedere la strategia di enumerazione che abbiamo utilizzata quando abbiamo enumerato Q^+ ed ottenere l'ordinamento definito ponendo

$$0 < 1 < 1/2 < 2/1 < 1/3 < 3/1 < 1/4 < 4/1 < 2/3 < 3/2 < \dots$$

Similmente, non è difficile trovare anche un buon ordinamento per l'intero insieme dei razionali Q intrecciando opportunamente razionali positivi e negativi.

6. Gruppi, anelli e campi

Ricordiamo brevemente alcune nozioni elementari di carattere algebrico. La più importante è forse quella di gruppo.

Definizione 6.1. Diciamo che una struttura algebrica $(D, \cdot, ^{-1}, 1)$ è un *gruppo* se:

- i) $(x \cdot y) \cdot z = x \cdot (y \cdot z)$ (proprietà associativa)
- ii) $x \cdot 1 = x$; $1 \cdot x = x$ (1 è elemento neutro)
- iii) $x \cdot x^{-1} = 1$; $x^{-1} \cdot x = 1$ (invertibilità)

Un gruppo è detto *commutativo* o *abeliano* se

- iv) $x \cdot y = y \cdot x$.

Definizione 6.2. Chiamiamo *anello unitario commutativo* ogni struttura algebrica $(D, +, \cdot, 0, 1)$ tale che:

- 1) $(D, +, 0)$ sia un gruppo commutativo
- 2) $(D, \cdot, 1)$ sia una operazione associativa e commutativa con 1 come elemento neutro
- 3) valga la proprietà distributiva del prodotto rispetto la somma, cioè

$$(a+b) \cdot c = a \cdot c + b \cdot c$$

Esempio di anello commutativo è quello degli interi relativi. Un altro esempio è quello degli interi modulo un fissato intero m .

Proposizione 6.3. In ogni anello risulta che:

- i) $x \cdot 0 = 0$.
- ii) $x \cdot (-y) = -x \cdot y$.
- iii) $x \cdot (-1)$ è l'opposto di x .
- iv) $(-1)^2 = 1$.

Dim. Per provare i) osserviamo che per la proprietà distributiva $x \cdot 0 = x \cdot (0+0) = x \cdot 0 + x \cdot 0$ da cui, sottraendo da entrambi i membri $x \cdot 0$, si ricava che $0 \cdot x = 0$. Per provare ii) osserviamo che

$$x \cdot (-y) + x \cdot y = x \cdot (-y + y) = x \cdot 0 = 0.$$

Le rimanenti proprietà sono ovvie. □

Definizione 6.4. Un anello unitario commutativo $(D, +, \cdot, 0, 1)$ è chiamato *campo* se $(D - \{0\}, \cdot, 1)$ è un gruppo.

Nella teoria degli anelli è importante la nozione di divisore dello zero.

Definizione 6.5. Prende il nome di *divisore dello zero* un elemento $x \neq 0$ tale che esiste $y \neq 0$ tale che $x \cdot y = 0$.

Proposizione 6.6. I divisori dello zero non sono invertibili. Pertanto in un campo non esistono divisori dello zero.

Dim. Sia x un divisore dello zero e supponiamo che esista l'inverso x^{-1} di x . Allora se $y \neq 0$ è tale che $x \cdot y = 0$ si avrebbe che $y = x^{-1} \cdot (x \cdot y) = x^{-1} \cdot 0 = 0$. □

Spesso in un campo viene definita una relazione d'ordine. In tale caso si richiede che tale relazione "si comporti bene" rispetto alle operazioni.

Definizione 6.7. Un anello unitario commutativo $(D, +, \cdot, 0, 1)$ è detto *ordinato* se è definita in D una relazione d'ordine totale $\leq$ compatibile con le operazioni cioè tale che

$$1) \quad a \leq b \Rightarrow a+c \leq b+c,$$

$$2) \quad a \leq b \Rightarrow ac \leq bc \text{ per ogni } c \geq 0.$$

Chiamiamo *positivi* gli elementi strettamente maggiori di 0 e *negativi* gli elementi strettamente minori di 0.

La struttura algebrica degli interi relativi è un esempio di anello ordinato. La struttura algebrica definita dai razionali è un tipico esempio di campo ordinato. Si dimostra che nell'anello degli interi modulo m e nell'anello dei numeri complessi non è possibile definire una relazione d'ordine che sia compatibile con le operazioni.

Proposizione 6.8. In ogni anello ordinato risulta che:

$$i) \quad a \leq b \Rightarrow a-b \leq 0 ; \quad ii) \quad a \leq b \Rightarrow 0 \leq b-a.$$

$$iii) \quad b \geq 0 \Leftrightarrow -b \leq 0 ; \quad iv) \quad c \geq 0, b \geq 0 \Rightarrow b \cdot c \geq 0$$

$$v) \quad c \geq 0, a \leq 0 \Rightarrow ac \leq 0 ; \quad vi) \quad c \leq 0, b \leq 0 \Rightarrow b \cdot c \geq 0$$

$$vii) \quad 1 \geq 0, -1 \leq 0.$$

Dim. L'implicazione *i)* si ottiene ponendo $c = -b$ in 1). La *ii)* si ottiene ponendo $c = -a$. Se in *i)* si pone $a = 0$ allora si ottiene $b \geq 0 \Rightarrow -b \leq 0$. Se in *ii)* si pone $b = 0$ si ottiene $a \leq 0 \Rightarrow 0 \leq -a$. In questo modo la *iii)* è dimostrata. Le rimanenti proprietà si dimostrano in modo analogo. $\square$

7. La nozione generale di struttura

In questo paragrafo cominciamo col dare qualche idea generale della nozione di struttura matematica che definiamo come un insieme D con delle operazioni, delle relazioni ed alcuni elementi che giocano un ruolo particolare (come ad esempio gli elementi neutri).

Definizione 7.1. Una *struttura del primo ordine* è un oggetto matematico del tipo $(D, h_1, \dots, h_n, \mathcal{R}_1, \dots, \mathcal{R}_m, c_1, \dots, c_k)$ con D in-

sieme non vuoto detto *dominio*, $h_1, \dots, h_n$ operazioni, $\mathcal{R}_1, \dots, \mathcal{R}_m$ relazioni e $c_1, \dots, c_k$ elementi di S .

Se non esistono relazioni allora la struttura prenderà il nome di *struttura algebrica*, se non esistono operazioni prenderà il nome di *struttura relazionale*. I gruppi e gli anelli sono esempi di strutture algebriche. Gli insiemi ordinati sono esempi di strutture relazionali. Gli anelli ordinati (come $\mathbb{Z}$) costituiscono un esempio di struttura algebrica in cui sono presenti sia una relazione che operazioni.

Definizione 7.2. Due strutture $(D, h_1, \dots, h_n, \mathcal{R}_1, \dots, \mathcal{R}_m, c_1, \dots, c_k)$ e $(D', h'_1, \dots, h'_n, \mathcal{R}'_1, \dots, \mathcal{R}'_m, c'_1, \dots, c'_k)$ si dicono *dello stesso tipo* se

- i) per ogni i , h_i ed h'_i hanno lo stesso numero di variabili
- ii) per ogni j , $\mathcal{R}_j$ e $\mathcal{R}'_j$ si applicano allo stesso numero di elementi.

Ad esempio gli anelli unitari ed i campi sono strutture dello stesso tipo poiché sono forniti di un prodotto, di una somma, di costanti 0 ed 1. I campi ordinati sono di tipo diverso dai campi poiché hanno anche una relazione d'ordine. Nel seguito se $\mathcal{R}$ è una relazione su un insieme D ed $X \subseteq D$, allora la restrizione di $\mathcal{R}$ ad X è la relazione

$$\mathcal{R} \cap X^n = \{(x_1, \dots, x_n) : x_1 \in X, \dots, x_n \in X, (x_1, \dots, x_n) \in \mathcal{R}\}.$$

Denotiamo tale restrizione con $\mathcal{R}/X$.

Definizione 7.3. Una *sottostruttura* di una struttura $S = (D, h_1, \dots, h_n, \mathcal{R}_1, \dots, \mathcal{R}_m, c_1, \dots, c_k)$ è una struttura $S' = (D', h'_1, \dots, h'_n, \mathcal{R}'_1, \dots, \mathcal{R}'_m, c'_1, \dots, c'_k)$, dello stesso tipo di S , tale che:

- i) $D' \subseteq D$
- ii) $h'_i = h_i/D'$
- iii) $\mathcal{R}'_i = \mathcal{R}_i/D'$
- iv) $c'_i = c_i$.

Pertanto una sottostruttura di S si ottiene fissando una parte D' di D che sia stabile rispetto alle operazioni $h_1, \dots, h_n$ e che contenga le costanti $c_1, \dots, c_k$.

Definizione 7.4. Siano

$$(D, h_1, \dots, h_n, \mathcal{R}_1, \dots, \mathcal{R}_m, c_1, \dots, c_k)$$

e

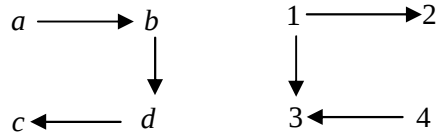
$$(D', h'_1, \dots, h'_n, \mathcal{R}'_1, \dots, \mathcal{R}'_m, c'_1, \dots, c'_k)$$

due strutture, chiamiamo *omomorfismo* una funzione $f : S_1 \rightarrow S_2$ tale che:

$$\begin{aligned} (x_1, \dots, x_n) \in \mathcal{R}_i &\Rightarrow (f(x_1), \dots, f(x_n)) \in \mathcal{R}'_i \\ f(h_i(x_1, \dots, x_n)) &= h'_i(f(x_1), \dots, f(x_n)) \\ f(c_i) &= c'_i. \end{aligned}$$

Diciamo che f è un *isomorfismo* se è un omomorfismo biiettivo il cui inverso è ancora un omomorfismo. Diciamo che f è una *immersione* se è un isomorfismo tra $(D, h_1, \dots, h_n, \mathcal{R}_1, \dots, \mathcal{R}_m, c_1, \dots, c_k)$ ed una sottostruttura di $(D', h'_1, \dots, h'_n, \mathcal{R}'_1, \dots, \mathcal{R}'_m, c'_1, \dots, c'_k)$.

Da notare che la nozione di isomorfismo per le strutture relazionali è leggermente diversa da quella data per le strutture algebriche in cui non si richiede che l'inverso f^{-1} sia ancora un omomorfismo perché ciò accade automaticamente. Invece per le strutture relazionali tale ipotesi è essenziale poiché esistono omomorfismi invertibili il cui inverso non è un omomorfismo. Ad esempio consideriamo due relazioni $\mathcal{R}_1$ ed $\mathcal{R}_2$ rappresentate dai due grafi



e sia f la funzione definita da $f(a) = 1, f(b) = 2, f(c) = 3, f(d) = 4$. Allora è immediato che f è un omomorfismo invertibile il cui inverso non è un omomorfismo in quanto $1\mathcal{R}_13$ ma non è vero che $f^{-1}(1)\mathcal{R}_2f^{-1}(3)$.

Si pone ora il problema se, data una relazione di equivalenza $\equiv$ in un insieme S in cui siano state definite relazioni ed operazioni se tali relazioni ed operazioni possono essere definite anche nel quoziente $S/\equiv$. Ad esempio, supponiamo che in S sia definita una operazione binaria $+$, allora ha senso, date due classi $X \in S/\equiv$ e $Y \in S/\equiv$, proporre il seguente algoritmo

- prendi un elemento $x \in X$
- prendi un elemento $y \in Y$

- calcola $x+y$
- considera la classe $[x+y]$.

Ma perché una tale definizione funzioni il risultato di un tale algoritmo non deve dipendere dal modo come x ed y sono scelti in X ed Y . In altre parole deve accadere che:

$$x \equiv x' \text{ e } y \equiv y' \Rightarrow x+y \equiv x'+y'.$$

Si perviene allora alla seguente definizione:

Definizione 7.5. Dato un insieme S ed una relazione di equivalenza $\equiv$, diciamo che una operazione n -aria $h : S^n \rightarrow S$ è *compatibile* con $\equiv$ se risulta:

$$x_1 \equiv x'_1, \dots, x_n \equiv x'_n \Rightarrow h(x_1, \dots, x_n) \equiv h(x'_1, \dots, x'_n)$$

Data una relazione n -aria $\mathcal{R}$, diciamo che $\equiv$ è *compatibile con* $\mathcal{R}$ se, per ogni x, x', y, y'

$$(x_1 \equiv x'_1, \dots, x_n \equiv x'_n) \Rightarrow (x_1, \dots, x_n) \in \mathcal{R} \Leftrightarrow (x'_1, \dots, x'_n) \in \mathcal{R}.$$

Data una struttura matematica $S = (D, h_1, \dots, h_n, \mathcal{R}_1, \dots, \mathcal{R}_m, c_1, \dots, c_k)$ chiamiamo *congruenza* una relazione di equivalenza $\equiv$ in D che sia compatibile sia con le operazioni che con le relazioni di tale struttura.

Definizione 7.6. Data una struttura matematica $S = (D, h_1, \dots, h_n, \mathcal{R}_1, \dots, \mathcal{R}_m, c_1, \dots, c_k)$ ed una congruenza $\equiv$, chiamiamo *quoziente di S modulo $\equiv$* , la struttura

$$S/\equiv = (D/\equiv, h'_1, \dots, h'_n, \mathcal{R}'_1, \dots, \mathcal{R}'_m, [c_1], \dots, [c_k])$$

dove si è posto

$$h'_i([x_1], \dots, [x_n]) = [h_i(x_1, \dots, x_n)],$$

$$\mathcal{R}'_i = \{([x_1], \dots, [x_n]) : (x_1, \dots, x_n) \in \mathcal{R}_i\}.$$

Esempio: Dato un intero m , si consideri in Z la relazione di *congruenza modulo m* , cioè la relazione definita dal porre $x \equiv y$ se e solo se $x-y$ è un multiplo di m . Allora $\equiv$ è una congruenza nell'anello Z . Il relativo quoziente viene detto *anello degli interi modulo m* .

Problema: Consideriamo la relazione $\equiv$ in Z definita dal considerare equivalenti due numeri che abbiano gli stessi divisori primi. Ad esempio avremo che $6 \equiv 18$ ma 6 non è equivalente a 730 e la classe $[6]$ contenente 6 è costituita da tutti i numeri che si possono costruire moltiplicando opportunamente 2 e 3 , $[6] = \{6,$

12, 18, ...}. Ancora, il numero 7 è equivalente solo ad una sua potenza e $[7] = \{7, 49, \dots\}$.

- la relazione $\equiv$ è di equivalenza ?
- la relazione $\equiv$ è compatibile con la moltiplicazione ?
- la relazione $\equiv$ è compatibile con l'addizione ?

Esempio. In $\mathbb{Z}$ consideriamo l'ordinamento usuale e sia $\equiv$ la congruenza modulo 5. In tali ipotesi $6 \equiv 1$, $5 \equiv 5$ ma pur essendo $1 < 5$ non è vero che $6 < 5$. Pertanto la congruenza modulo 5 non è compatibile con l'ordinamento in $\mathbb{Z}$. Ciò significa che $\equiv$ non può essere considerata una congruenza in $\mathbb{Z}$ inteso come anello ordinato.

Esercizio. Consideriamo nell'insieme $\{1, 2, 3, 4, 5, 6, 7, 8, 9, 10\}$ dei primi dieci numeri naturali, in cui è definita la solita relazione d'ordine le seguenti partizioni

$$P1 = \{\{1, 2, 10\}, \{3, 4, 5, 6, 9\}, \{7, 8\}\}.$$

$$P2 = \{\{1, 2, 3, 4\}, \{5, 6, 7\}, \{8, 9, 10\}\}$$

Tali partizioni determinano a loro volta due relazioni di equivalenza. Dire quale delle due è una congruenza e perché.

8. Sistemi di chiusura, operatori e punti fissi

Moltissime nozioni matematiche possono essere introdotte in termini di sistemi di chiusura e di operatori di chiusura.

Definizione 8.1. Sia S un insieme non vuoto e denotiamo con $\mathcal{P}(S)$ l'insieme delle sue parti, allora chiameremo *sistema di chiusura* una classe C di sottoinsiemi di S tale che

$$i) S \in C$$

ii) C è chiusa rispetto alle intersezioni, cioè l'intersezione di una qualunque famiglia di elementi di C appartiene ancora a C .

Dato un sistema di chiusura è definita la nozione di sottoinsieme generato da un dato insieme.

Definizione 8.2. Dato un sistema di chiusura C in un insieme non vuoto S ed un sottoinsieme X di S poniamo

$$\langle X \rangle = \bigcap \{Y \in C : Y \supseteq X\}$$

è diciamo che $\langle X \rangle$ è l'elemento di C generato da X .

Poiché $\langle X \rangle$ appartiene a C , è evidente che $\langle X \rangle$ è il più piccolo elemento di C che contiene X .

Esempio di carattere topologico. Ad esempio se C è la classe degli insiemi chiusi di R allora C è un sistema di chiusura in quanto l'intersezione di una famiglia di insiemi chiusi è ancora un insieme chiuso. Se X è un sottoinsieme di R allora $\langle X \rangle = \overline{X}$ cioè $\langle X \rangle$ è la chiusura dell'insieme X .

Esempio di carattere algebrico. Sia C la classe di tutti i sottogruppi di un dato gruppo G . Allora C è un sistema di chiusura e, per ogni sottoinsieme X di G , $\langle X \rangle$ è il sottogruppo generato da X .

Esempio di carattere geometrico. Sia C la classe di tutti gli insiemi convessi del piano euclideo. Allora C è un sistema di chiusura. Per ogni insieme X di punti, $\langle X \rangle$ è la *chiusura convessa* di X .

La nozione di sistema di chiusura è strettamente legata alla nozione di operatore di chiusura.

Definizione 8.3. Chiamiamo *operatore di chiusura* una funzione $T: \mathcal{P}(S) \rightarrow \mathcal{P}(S)$ tale che:

- a) $T(X) \supseteq X$ (proprietà di inclusione)
- b) $X \subseteq Y \Rightarrow T(X) \subseteq T(Y)$ (crescenza o monotonia)
- c) $T(T(X)) = T(X)$. (idempotenza).

Esempio di carattere topologico. Sia X un sottoinsieme di R ed indichiamo con $T(X)$ la chiusura $\overline{X}$ dell'insieme X . Allora T è un operatore di chiusura.

Esempio di carattere algebrico. In uno spazio vettoriale per ogni insieme X di vettori indichiamo con $T(X)$ l'insieme dei vettori linearmente dipendenti da X . Allora T è un operatore di chiusura.

Definizione 8.4. Dato un operatore T , si chiama *punto fisso* o *punto unito* di T un insieme X tale che $T(X) = X$.

Si osservi che la proprietà di idempotenza $T(T(X)) = T(X)$ equivale a dire che per ogni X l'insieme $T(X)$ è un punto fisso di T . Que-

sto significa che se T è un operatore di chiusura allora tutti gli elementi del condominio di T sono punti fissi.

La seguente proposizione mostra che gli operatori di chiusura sono strettamente collegati ai sistemi di chiusura.

Teorema 8.5. Sia $T : \mathcal{P}(S) \rightarrow \mathcal{P}(S)$ un operatore che verifica la proprietà di inclusione e quella di crescita. Allora l'insieme

$$C_T = \{X \in \mathcal{P}(S) : T(X) = X\}$$

dei punti uniti di T è un sistema di chiusura.

Dim. Sia $(X_i)_{i \in I}$ una famiglia di elementi di C_T , allora, essendo $\bigcap_{i \in I} X_i \subseteq X_j$ per ogni $j \in I$, per la crescita di T , risulta $T(\bigcap_{i \in I} X_i) \subseteq T(X_j)$. Ne segue che, essendo X_j punto fisso per T , $T(\bigcap_{i \in I} X_i) \subseteq X_j$ per ogni $j \in I$ e quindi $T(\bigcap_{i \in I} X_i) \subseteq \bigcap_{j \in I} X_j$. Poiché per la proprietà di inclusione $T(\bigcap_{i \in I} X_i) \supseteq \bigcap_{j \in I} X_j$, l'insieme $\bigcap_{i \in I} X_i$ è un punto fisso e quindi appartiene a C_T .

Proposizione 8.6. Supponiamo che T verifichi le proprietà di inclusione e di crescita e poniamo

$$D(X) = \bigcap \{Y : Y \supseteq X \text{ e } Y = T(Y)\}, \quad (8.1)$$

Allora $D(X)$ è il minimo punto fisso di T contenente X .

Dim. Per definizione di intersezione $D(X)$ è l'estremo inferiore dell'insieme $\{Y : Y \supseteq X \text{ e } Y = T(Y)\}$ dei punti fissi contenenti X . Poiché il teorema 7.5 ci dice che $D(X)$ appartiene a tale insieme, $D(X)$ risulta esserne il minimo. D'altra parte se Y è un punto fisso che contiene X , allora appartenendo all'insieme $\{Y : Y \supseteq X \text{ e } Y = T(Y)\}$ risulta contenere $D(X)$.

Teorema 8.7. Se T è un operatore di chiusura allora l'insieme C_T dei suoi punti fissi è un sistema di chiusura. Viceversa, sia C un sistema di chiusura e definiamo l'operatore $T_C : \mathcal{P}(S) \rightarrow \mathcal{P}(S)$ ponendo

$$T_C(X) = \langle X \rangle = \bigcap \{Y \in C : Y \supseteq X\}.$$

Allora T_C è un operatore di chiusura il cui insieme di punti fissi è C .

Dim. La prima parte del teorema segue dalla proposizione precedente in quanto un operatore di chiusura verifica la proprietà di inclusione ed è monotono. La seconda parte è evidente. Infatti è

immediato che J_C verifica le proprietà di inclusione e di monotonia. D'altra parte, se si osserva che

$$Y \supseteq T_C(X) \Leftrightarrow Y \supseteq X,$$

risulta anche

$$T_C(T_C(X)) = \cap \{Y \in C : Y \supseteq T_C(X)\} = \cap \{Y \in C : Y \supseteq X\}.$$

9. Due teoremi di punto fisso per operatori

La seguente proposizione mostra che gli operatori crescenti ammettono sempre punto fisso.

Teorema 9.1. Sia $T : \mathcal{P}(S) \rightarrow \mathcal{P}(S)$ un operatore crescente. Allora l'insieme

$$Z = \cup \{X : X \subseteq T(X)\}$$

è un punto fisso di T .

Dim. Infatti per definizione per ogni $X \subseteq T(X)$ risulta che $X \subseteq Z$ e quindi, essendo T crescente, $X \subseteq T(X) \subseteq T(Z)$. Ne segue che $Z = \cup \{X : X \subseteq T(X)\} \subseteq T(Z)$. Da tale disuguaglianza segue che $T(Z) \subseteq T(T(Z))$ e che quindi $T(Z)$ appartiene all'insieme $\{X : X \subseteq T(X)\}$. Conseguentemente $T(Z) \subseteq Z$. Avendo già osservato che $Z \subseteq T(Z)$, possiamo concludere che $T(Z) = Z$.

Nel caso di operatori algebrici esiste un modo più “costruttivo” per ottenere un punto fisso.

Definizione 9.2. Chiamiamo *operatore algebrico* su S ogni funzione $T : \mathcal{P}(S) \rightarrow \mathcal{P}(S)$ tale che:

- a) $T(X) \supseteq X$ (proprietà di inclusione)
 - b) $X \subseteq Y \Rightarrow T(X) \subseteq T(Y)$ (crescenza o monotonia)
 - c) $x \in T(X) \Rightarrow \exists F \subseteq X, F \text{ finito}, x \in T(F)$ (compattezza o finitezza).
- T è un *operatore di chiusura algebrico*, cioè se oltre ad essere un operatore algebrico verifica anche
- d) $T(T(X)) = T(X)$ (proprietà di chiusura).

La teoria degli operatori algebrici deve essere vista come un modo astratto di considerare un processo con cui si costruiscono nuovi oggetti a partire da un dato insieme X di oggetti. Allora

a) significa che tra le cose che posso costruire con X ci sono gli elementi di X ,

b) significa che se una cosa può essere costruita a partire da X allora (a maggior ragione) può essere costruita a partire da un insieme che contiene X ,

c) afferma che dire che una cosa è costruibile a partire da X significa in realtà che è costruibile a partire da un numero finito di elementi di X . In altre parole, significa che il processo di costruzione è finitario.

d) si può interpretare dicendo che se un oggetto x è costruito con materiale in $T(X)$ che a sua volta è stato costruito con materiale in X allora tale oggetto è, di fatto, costruito con materiale in X . In altre parole a partire dagli oggetti in $T(X)$ non è possibile costruire niente che non sia già in $T(X)$.

Da notare che le condizioni b) e c) equivalgono a dire che

$$T(X) = \cup \{T(F) : F \subseteq X \text{ e } F \text{ finito}\}$$

cioè equivalgono a dire che il calcolo di $T(X)$ si può effettuare riferendosi solo alla parti finite di T .

Problema. Sia S un insieme e Z un suo sottoinsieme. Dire quali proprietà verifica l'operatore T definito ponendo, per ogni sottoinsieme X di S

$$T(X) = X \cup Z.$$

Problema. Sia T l'operatore che associa ad ogni insieme X di numeri interi l'insieme $T(X)$ dei divisori degli elementi di X . Ad esempio $T(\{3,14,15\}) = \{1,3,2,7,5,14,15\}$. Dire che tipo di operatore è.

Per caratterizzare gli operatori algebrici è utile il seguente lemma.

Lemma 9.3. Sia $(X_n)_{n \in \mathbb{N}}$ una successione di sottoinsiemi di S crescente rispetto alla inclusione ed F un sottoinsieme finito di S , allora

$$F \subseteq \cup_{n \in \mathbb{N}} X_n \Rightarrow \exists j \in \mathbb{N} \text{ tale che } F \subseteq X_j. \quad (9.1)$$

Dim. Procederemo per induzione sulla cardinalità $n = |F|$ di F . Per $n = 1$ l'implicazione (8.1) è conseguenza immediata della definizione di unione. Supponiamo che (8.1) sia vera per un insieme di n elementi, sia F un insieme di $n+1$ elementi e supponiamo che $F \subseteq \cup_{n \in \mathbb{N}} X_n$. Allora sarà $F = F' \cup \{a\}$ con a opportuno elemento e $|F'| = n$. Pertanto, per ipotesi di induzione, essendo

$F' \subseteq F \subseteq \bigcup_{n \in \mathbb{N}} X_n$ possiamo ricavare l'esistenza di un intero h tale che $F' \subseteq X_h$. Sia k tale che $a \in X_k$ e sia $j = \max\{h, k\}$, allora dalla crescenza di $(X_n)_{n \in \mathbb{N}}$ si ricava che $F \subseteq X_h \cup X_k = X_j$.

Problema. Mostrare che il lemma ora enunciato non vale senza l'ipotesi di crescenza per $(X_n)_{n \in \mathbb{N}}$ e che non vale senza l'ipotesi di finitezza per F .

Lemma 9.4. Se un operatore T è compatto e monotono allora per ogni successione crescente $(X_n)_{n \in \mathbb{N}}$ di sottoinsiemi di S risulta che

$$T(\bigcup_{n \in \mathbb{N}} X_n) = \bigcup_{n \in \mathbb{N}} T(X_n). \quad (9.2)$$

Dim. Per ogni fissato m , $X_m \subseteq \bigcup_{n \in \mathbb{N}} X_n$ e quindi per la monotonia $T(X_m) \subseteq T(\bigcup_{n \in \mathbb{N}} X_n)$. Ne segue che

$$\bigcup_{n \in \mathbb{N}} T(X_n) = \bigcup_{m \in \mathbb{N}} T(X_m) \subseteq T(\bigcup_{n \in \mathbb{N}} X_n).$$

Per provare l'inclusione inversa, sia $x \in T(\bigcup_{n \in \mathbb{N}} X_n)$, allora per la compattezza esiste un sottoinsieme finito F di $\bigcup_{n \in \mathbb{N}} X_n$ tale che $x \in T(F)$. Detto j un intero tale che $F \subseteq X_j$ risulta anche che $x \in T(X_j)$ e quindi che $x \in \bigcup_{n \in \mathbb{N}} T(X_n)$. Pertanto $T(\bigcup_{n \in \mathbb{N}} X_n) \subseteq \bigcup_{n \in \mathbb{N}} T(X_n)$.

Il seguente teorema mostra che ogni operatore compatto e monotono ammette un punto fisso. Nel seguito scriveremo $T^n(X)$ per indicare il risultato della applicazione n volte dell'operatore T ad X ; più precisamente definiamo $T^n(X)$ per ricorsione su n tramite le equazioni

- $T^0(X) = X$
- $T^{n+1}(X) = T(T^n(X))$.

Teorema 9.5. Sia $T : \mathcal{P}(S) \rightarrow \mathcal{P}(S)$ un operatore che verifica (8.2) per ogni successione crescente $(X_n)_{n \in \mathbb{N}}$ di sottoinsiemi di S . Allora se X è un insieme tale che $T(X) \supseteq X$, l'insieme $\bigcup_{n \in \mathbb{N}} T^n(X)$ è il minimo punto fisso di T contenente X . In particolare, $\bigcup_{n \in \mathbb{N}} T^n(\emptyset)$ è il minimo punto fisso di T .

Dim. La condizione (9.2) comporta la crescenza di T . Pertanto dall'ipotesi $T(X) \supseteq X$, applicando n volte l'operatore T , segue che $T^{n+1}(X) \supseteq T^n(X)$. Essendo quindi $(T^n(X))_{n \in \mathbb{N}}$ una successione crescente per 8.2)

$$T(\bigcup_{n \in \mathbb{N}} T^n(X)) = \bigcup_{n \in \mathbb{N}} T^{n+1}(X) = \bigcup_{n \in \mathbb{N}} T^n(X).$$

Ciò prova che $\cup_{n \in \mathbb{N}} T^n(X)$ è un punto fisso di T . Per provare che tale insieme è il minimo punto fisso contenente X , sia M un qualsiasi punto fisso contenente X . Allora per la monotonia di T , applicando n volte T alla disuguaglianza $M \supseteq X$, otteniamo che $T^n(M) \supseteq T^n(X)$ e quindi, essendo M punto fisso, che $M \supseteq T^n(X)$. In definitiva $M \supseteq \cup_{n \in \mathbb{N}} T^n(X)$ e quindi $\cup_{n \in \mathbb{N}} T^n(X)$ è il minimo punto fisso contenente X .

Corollario 9.6. Sia $T : \mathcal{P}(S) \rightarrow \mathcal{P}(S)$ un operatore algebrico, allora per ogni $X \subseteq S$ l'equazione

$$D(X) = \cup_{n \in \mathbb{N}} T^n(X) \quad (9.3)$$

fornisce il minimo punto fisso $D(X)$ di T contenente X .

Esempio: Sia S uno spazio Euclideo e definiamo l'operatore $T : \mathcal{P}(S) \rightarrow \mathcal{P}(S)$ ponendo

$$T(X) = \{x \in S \mid \exists p, \exists q \in X, x \in \underline{pq}\}$$

avendo indicato con $\underline{pq}$ il segmento chiuso di estremi p e q . In altre parole $T(X)$ è l'insieme dei punti che si trovano su di un segmento i cui estremi appartengono a X . Sono punti uniti di T tutti e soli i sottoinsiemi convessi di S . T è un operatore algebrico di tipo due. Infatti $x \in T(X)$ implica che $x \in \underline{pq}$ con $p \in X$ e $q \in X$, e quindi $x \in T(\{p, q\})$. Ne segue che:

- l'intersezione di una famiglia di insiemi convessi è ancora un insieme convesso;
- dato un insieme X esiste il più piccolo insieme convesso $D(X)$ contenente X e lo si può ottenere come unione della catena $T^1(X)$, $T^2(X)$, $\dots$

10. Come generare relazioni di ordine o di equivalenza

Se una relazione $\mathcal{R}$ non è una relazione di pre-ordine allora non è difficile farla diventare una relazione di pre-ordine aggiungendo un opportuno insieme di coppie. Per mostrare come questo può essere fatto vediamo prima come si può rendere riflessiva una relazione.

Definizione 10.1. Dato un insieme S , la *diagonale* di S è la relazione

$$D(S) = \{(x, y) : x \text{ coincide con } y\}.$$

Tale relazione viene anche chiamata *identità*.

La relazione $D(S)$ viene anche detta *diagonale di S* in quanto nel caso $S = R$, la sua rappresentazione in un sistema di assi cartesiani ortogonali coincide con la diagonale del primo e terzo quadrante.

Dalla definizione di relazione riflessiva segue immediatamente la seguente proposizione.

Definizione 10.2. Dato un insieme non vuoto S indichiamo con $Rifl : \mathcal{P}(S \times S) \rightarrow \mathcal{P}(S \times S)$ l'operatore che associa ad ogni relazione binaria $\mathcal{R}$, la relazione

$$Rifl(\mathcal{R}) = \mathcal{R} \cup D(S).$$

Vale la seguente proposizione di cui omettiamo la semplice dimostrazione.

Proposizione 10.3. $Rifl$ è un operatore di chiusura, inoltre le seguenti asserzioni sono equivalenti:

- i) $\mathcal{R}$ è riflessiva
- ii) $Rifl(\mathcal{R}) = \mathcal{R}$ cioè $\mathcal{R}$ è un punto fisso di $Rifl$
- iii) $\mathcal{R} \supseteq D(S)$.

Inoltre $Rifl(\mathcal{R})$ è la più piccola relazione riflessiva contenente $\mathcal{R}$.

La relazione $Rifl(\mathcal{R})$ viene chiamata *relazione riflessiva generata da $\mathcal{R}$* . Vediamo ora come si può rendere transitiva una relazione.

Definizione 10.4. Sia $\mathcal{R}$ una relazione in S , chiamiamo *percorso* da x ad y una successione $x_1, \dots, x_n$ di elementi di S tale che $x = x_1$, $x_n = y$ e $x_{i-1} \mathcal{R} x_i$. Definiamo l'operatore $Trans : \mathcal{P}(S \times S) \rightarrow \mathcal{P}(S \times S)$ che associa ad ogni relazione binaria $\mathcal{R}$, la relazione

$$Trans(\mathcal{R}) = \{(x, y) : \text{esiste un percorso da } x \text{ ad } y\}.$$

Vale la seguente proposizione.

Proposizione 10.5. $Trans$ è un operatore di chiusura. Una relazione binaria $\mathcal{R}$ è transitiva se e solo se $\mathcal{R} = Trans(\mathcal{R})$ cioè se e solo se è un punto fisso di $Trans$. Inoltre $Trans(\mathcal{R})$ è la più piccola relazione transitiva contenente $\mathcal{R}$.

Dim. E' evidente che $\mathcal{R}$ è transitiva se e solo se contiene $\text{Trans}(\mathcal{R})$. Per provare che $\mathcal{R}' = \text{Trans}(\mathcal{R})$ è transitiva supponiamo che $x\mathcal{R}'y$ e $y\mathcal{R}'z$ e quindi che esista un percorso $a_1, \dots, a_n$ da x ad y ed un percorso $b_1, \dots, b_m$ da y a z . Allora $a_1, \dots, a_n, b_1, \dots, b_m$ è un percorso da x a z . Per provare che $\mathcal{R}'$ è la più piccola relazione transitiva contenente $\mathcal{R}$ indichiamo con $\mathcal{R}^*$ una relazione transitiva contenente $\mathcal{R}$. Allora $\mathcal{R}^* \supseteq \text{Trans}(\mathcal{R}^*) \supseteq \text{Trans}(\mathcal{R})$.

Proposizione 10.6. La relazione $\text{Trans}(\text{Rifl}(\mathcal{R}))$ è sia riflessiva che transitiva e quindi è una relazione di pre-ordine. Precisamente è il più piccolo preordine contenente $\mathcal{R}$ e viene chiamato *pre-ordine generato da $\mathcal{R}$* .⁴

Resta da vedere come sia possibile trasformare una relazione di pre-ordine in una relazione d'ordine, cioè ottenere la proprietà antisimmetrica. Ora se $\mathcal{R}$ contiene due coppie (a, b) e (b, a) con $a \neq b$, allora non esiste nessuna speranza di estendere $\mathcal{R}$ in modo da ottenere tale proprietà. Per ottenerla allora dobbiamo quotizzare opportunamente la struttura $(S, \mathcal{R})$. Infatti vale la seguente proposizione.

Proposizione 10.7. Sia $\leq$ una relazione di un pre-ordine, allora la relazione $\equiv$ che si ottiene ponendo $x \equiv y$ se e solo se $x \leq y$ e $y \leq x$ è una relazione di equivalenza. Se si definisce in $S/\equiv$ la relazione $\leq$ ponendo

$$[x] \leq [y] \Leftrightarrow x \leq y$$

tale relazione è d'ordine in $S/\equiv$.

Se invece di mirare a relazioni d'ordine vogliamo ottenere relazioni di equivalenza, dobbiamo vedere come si possa modificare una relazione $\mathcal{R}$ in modo da ottenere la simmetria.

⁴Abbiamo già utilizzato questo modo di procedere quando abbiamo definito la relazione d'ordine in una terna di Peano come la relazione generata dalla relazione "successivo". In questo caso un percorso da x ad y è costituito da una successione del tipo $s^0(x), s^1(x), \dots, s^n(x) = y$.

Proposizione 10.8. Sia $Simm : \mathcal{P}(S \times S) \rightarrow \mathcal{P}(S \times S)$ l'operatore che associa ad ogni relazione binaria $\mathcal{R}$, la relazione

$$Simm(\mathcal{R}) = \mathcal{R} \cup \mathcal{R}^{-1}.$$

Allora $Simm$ è un operatore di chiusura. Inoltre le seguenti asserzioni sono equivalenti:

- i) $\mathcal{R}$ è simmetrica
- ii) $\mathcal{R}$ è un punto fisso di $Simm$
- iii) $\mathcal{R} \supseteq \mathcal{R}^{-1}$ è un punto fisso di Si .

Ne segue che $Simm(\mathcal{R})$ è la più piccola relazione simmetrica contenente $\mathcal{R}$.

$Simm(\mathcal{R})$ viene chiamata *relazione simmetrica generata da $\mathcal{R}$* .

Proposizione 10.9. Definiamo l'operatore $Eq : \mathcal{P}(S \times S) \rightarrow \mathcal{P}(S \times S)$ ponendo

$$Eq(\mathcal{R}) = Trans(Simm(Rifl(\mathcal{R}))).$$

Allora Eq è un operatore di chiusura i cui punti fissi coincidono con le relazioni di equivalenza. Ne segue che $Eq(\mathcal{R})$ è la più piccola relazione di equivalenza contenente $\mathcal{R}$ e viene chiamata *relazione di equivalenza generata da $\mathcal{R}$* .

INDICE ANALITICO

addizione di due cardinali; 179
alfabeto; 249
algebra dei numeri cardinali; 180
algebra dei numeri ordinali; 185
anello degli interi relativi; 125
anello ordinato; 295
anello unitario commutativo; 293
archimedeo; 216
Aristotele
 sillogismi; 273
Aristotele, quadrilatero; 271
aritmetizzazione; 105
Assioma del singoletto; 222
Assioma dell'estensionalità; 221
Assioma dell'infinito; 224
Assioma dell'insieme delle parti; 224
Assioma dell'insieme vuoto; 222
Assioma dell'unione; 223
Assioma dell'unione di due insiemi; 223
Assioma della scelta; 225
Assioma di continuità; 27
Assioma di isolamento; 219; 224
Assioma di sostituzione; 225
assiomi logici; 259
biettiva; 282
Brouwer L. E. J.; 207
buon ordinamento, Principio; 227
buon ordine come proprietà del secondo ordine; 257
buon ordine in una terna di Peano; 117
buon ordine, relazione; 288
campo; 294
campo completo; 217
campo dei numeri razionali; 127
campo dei numeri reali
 definizione assiomatica; 217
campo dei numeri reali tramite le sezioni di Dedekind; 128
campo dei numeri reali tramite le successioni di Cauchy; 131
campo dei razionali; 125
campo dei razionali non-standard; 140
campo dei reali non standard; 142
Cartesio; 81
 Discorso sul Metodo; 86
 La Geometria; 82

classe completa di equivalenza; 285
congruenza; 298
coppia; 281
costruzione delle equazioni; 89
definizioni per astrazione; 284
diagonale; 305; 306
dimostrazione in logica matematica; 260
dimostrazione per assurdo; 11
elevazione a potenza tra cardinali; 181
equicompletabilità; 32
equiconvergenti; 133
equipotenti; 158
equiscomponibilità; 32
 teorema; 33
 teorema inverso; 35
estremo inferiore; 287
estremo superiore; 287
Euclide
 Elementi; 20
falsa confutazione dell'assioma della scelta; 228
false equiscomponibilità; 38
falsi teoremi euclidei; 93
falso uso del principio di induzione; 109
filtro; 136
filtro principale; 137
forma normale; 149
formule ben formate; 258
Frege; 206
funzione computabile; 255
generalizzazione, regola; 259
geometrie finite; 213
geometrie non euclidee; 73
grandezze omogenee; 25
gruppo; 293
Hilbert
 Fondamenti della Geometria"; 210
 Sull'infinito; 244
idealizzazione degli enti matematici; 17
immersione; 297
implicatura; 273
infinitamente vicini; 216
infinitesimo elemento; 215
infinito elemento; 215
infinito insieme; 160
iniettiva; 282

insieme decidibile; 254
insieme fuzzy; 190
insieme infinito; 120
insieme rozzo; 191
insieme vago; 190
insiemi enumerabili; 161
insiemi numerabili; 161
interpretazione; 261
intuizionismo; 207
Ipotesi del continuo; 229
Ipotesi generalizzata del continuo; 229
Lemma di Zorn; 227
linguaggio dei termini; 256
linguaggio del primo ordine; 256; 257
linguaggio del secondo ordine; 257
linguaggio formale; 250
logica del primo ordine; 257
logica del secondo ordine; 257
maggiorante; 287
massimo; 287
Mathematica; 149
metodo di diagonalizzazione di Cantor; 168
minimo; 287
minorante; 287
modelli di geometrie non euclidee; 77
modello di Klein; 77
modello di Poincaré; 79
Modus Ponens; 259
moltiplicazione di cardinali; 179
multinsieme; 189
notazione infissa; 258
notazione prefissa; 257
nozioni comuni; 21
numeri algebrici; 254
numeri cardinali; 177
numeri complessi; 176
numeri naturali; 118
 nella scuola Pitagorica; 1
numeri ordinali; 183
numeri razionali; 125
numeri razionali non-standard; 140
numeri reali
 sistema di assiomi; 215
numeri reali non-standard; 142
numeri relativi; 122

numeri trascendenti; 254
numero di codice; 269
numero naturale non-standard; 268
omomorfismo; 297
operatore di chiusura; 300
paradosso "io sono una asserzione falsa"; 268
paradosso degli otto raggi; 38
paradosso dei quadrati perfetti; 154
paradosso del barbiere; 206
paradosso del mucchio di grano; 110; 118
paradosso del quadrato e del cerchio; 314
paradosso della classe degli insiemi con tre elementi; 204
paradosso della classe degli insiemi finiti; 204
paradosso della classe dei gruppi; 204
paradosso della classe di tutti gli insiemi; 202
paradosso della incommensurabilità tramite interi; 8
paradosso della incommensurabilità tramite razionali; 9
paradosso della scomposizione del quadrato; 96
paradosso della somma di infiniti 1 ed -1; 155
paradosso delle due monete; 155
paradosso delle ruote concentriche; 154
 soluzione; 173
paradosso di Achille e la tartaruga; 15
paradosso di Banach-Tarski o della duplicazione dei pani e dei pesci; 38
paradosso di Berry; 118
paradosso di Russell; 203
parola; 249
parte stabile; 120
parte stabile generata; 121
Peano; 210
Peano curva; 175
Pitagora; 1
postulati; 22
Postulato di Archimede; 26
potenza del continuo; 168; 172
pre-ordine generato; 307
Principio di comprensione; 201
Principio di induzione; 108
Principio di induzione transfinita; 290
Principio di sostanzialità; 201
problema di Hilbert terzo; 35; 37
prodotto cartesiano; 281
prodotto di due insiemi ben ordinati; 184
proiezione i -esima; 263
punto fisso di un operatore; 300

quoziente; 285
razionali non-standard; 135
regolare insieme; 103
relazione binaria; 281
relazione d'ordine; 286
relazione di buon ordine; 288
relazione di equivalenza; 284
relazione di equivalenza generata; 308
relazione di pre-ordine; 286
relazione riflessiva generata; 306
relazione simmetrica generata; 308
reticolo come insieme ordinato; 287
reticolo come struttura algebrica; 288
reticolo completo; 287
ricorsione; 110; 111; 114
Sesto Empirico; 40
sezione del campo dei numeri razionali; 129
sillogismo; 273
 barbara, darii, celarent, ferio, barbari, celaront; 276
simile; 82
sistema di chiusura; 299
sistema di numeri naturali; 118
somma di due insiemi ordinati; 185
sottostruttura; 296
spazio metrico; 60
spazio pseudo-metrico; 60
stoici, contributo alla logica; 276
struttura algebrica; 296
struttura del primo ordine; 295
struttura relazionale; 296
successione di Cauchy; 133
successivo in una insieme ben ordinato; 288
successivo in una terna di Peano; 106
suriettiva; 282
Teorema dei numeri primi; 12; 13
Teorema dell'angolo esterno; 73
Teorema di Cantor; 177
Teorema di Cantor-Bernstein; 308
Teorema di Cohen; 229
Teorema di Completezza; 266
Teorema di Euclide; 85
Teorema di Gödel (primo); 268
Teorema di Gödel (secondo); 270
Teorema di Pitagora; 5
Teorema di punto fisso per operatori compatti e monotoni; 304

Teorema fondamentale dell'aritmetica; 291
Teorema inverso di Pitagora; 7
teoria categorica; 231
teoria completa; 232
teoria consistente; 231
teoria delle classi; 218
teoria delle proporzioni; 29
teoria di Zermelo-Fraenkel; 220
teoria indipendente; 232
teoria soddisfacibile; 231
Terne di Peano, teoria del primo ordine; 267
Terne di Peano, teoria del secondo ordine; 105
ultrafiltro; 139
unione disgiunta; 180
verità di una formula; 266
Whitehead A. N.; 56

BIBLIOGRAFIA.

Per la storia della matematica

- Morris Kline, *La matematica nella cultura occidentale*, Feltrinelli.
- Morris Kline, *Storia del pensiero matematico*, vol I e II, Einaudi, 1972.
- L. L. Radice, *L'infinito*, Editori Riuniti.
- Bottazzini-Freguglia-Rigatelli, *Fonti per la storia della matematica*, Sansoni, 1992.
- B. D'Amore, M. Matteuzzi, *Gli interessi matematici*, Marsilio.
- G. Lolli, *Da Euclide a Goedel*, Il Mulino, 2004.
- Biacino Loredana: *Le funzioni elementari: un approccio storico*, Ed. CompoMat, 2009.

Per quanto riguarda i fondamenti della geometria.

- D. Hilbert, *Fondamenti della geometria*, Feltrinelli, 1970.
- P. Odifreddi, *Divertimento geometrico: Le origini geometriche della logica da Euclide a Hilbert*, Boringhieri, 2003.
- E. Agazzi, D. Palladino, *Le geometrie non euclidee*, Mondadori.
- R. Trudeau, *La rivoluzione euclidea*, Boringhieri, 1991.

Per chi ha interessi verso la filosofia della matematica.

- C. Cellucci, *La filosofia della matematica del Novecento*, Laterza, 2007.
- E. Casari, *La filosofia della matematica del '900*, Sansoni.
- E. Casari, *Questioni di filosofia della matematica*, Feltrinelli.
- G. Lolli, *Filosofia della matematica: L'eredità del Novecento*
- L. Geymonat, *Storia del pensiero filosofico e scientifico*, Garzanti.
- Rudy Rucker, *La mente e l'infinito*, Muzzio, 1991.
- Wang Hao, *Dalla matematica alla filosofia*, Boringhieri, 1984.
- D. R. Hofstadter, Goedel, Escher, *Bach: un'eterna Ghirlanda Brillante*. Il Mulino, 2002.

Per chi ha interesse al rapporto tra matematica e letteratura
C. Toffalori, *L'aritmetica di cupido*, Guanda Editore 2011.

Per una semplice e rigorosa introduzione dei sistemi numerici.

- M. R. Enea, D. Saeli, *Sistemi Numerici*, ARACNE editrice 2009.