

FISICA MATEMATICA I

LEZIONI DI SANDRO GRAFFI

**I anno del Corso di studi triennale in matematica presso
l'Università di Bologna, I semestre dell'Anno accademico 2001/2002**

Premessa

Nel presente anno accademico 2000/2001 hanno avuto inizio i corsi previsti dal nuovo ordinamento degli studi universitari, noto a tutti come "3+2". Le caratteristiche fondamentali del nuovo ordinamento sono queste: si comincia con un ciclo di istruzione triennale che permette di conseguire un titolo di studio (Laurea breve); il laureato che ne è in possesso può poi proseguire gli studi per un ulteriore ciclo biennale al termini del quale consegue la cosiddetta Laurea specialistica.

Il risultato principale atteso dai promotori del riordino è la diminuzione del fenomeno dell'abbandono degli studi prima del conseguimento della laurea; a questo scopo si vuole che già la laurea breve dia agli studenti una preparazione adeguata all'immissione nel mondo del lavoro. Ne consegue che il riordino necessariamente comporta una ristrutturazione profonda dei contenuti dei tradizionali corsi di laurea quadriennali o quinquennali. Le scorciatoie tipo compressione in tre anni di quanto si faceva in quattro, oppure cancellazione pura e semplice dell'ultimo anno, comprensibili data la fretta dimostrata dal governo nel portare a termine una simile rivoluzione, avranno a mio parere vita breve.

La ristrutturazione si presenta particolarmente delicata nei corsi di Laurea scientifici "duri" (Matematica, Fisica, Ingegneria) perché essa porta con sé lo smantellamento del glorioso biennio, pilastro e vanto della formazione nelle scienze esatte nel nostro paese. I tradizionali corsi annuali di Analisi Matematica, Geometria ecc. vengono progressivamente sostituiti da vari "moduli" semestrali strutturati in modo non dissimile dai corsi di "Calculus" o "Linear Algebra" ecc. professati nelle Università americane. Un simile alleggerimento aggrava ulteriormente la difficoltà di toccare nell'insegnamento prelaurea almeno alcuni degli sviluppi recenti della matematica. D'altra parte, questo è un aspetto che assume evidentemente un carattere fondamentale nel Corso di Laurea dedicato proprio alla Matematica.

Il Consiglio di Corso di Laurea in Matematica dell'Università di Bologna ha ritenuto pertanto di introdurre qualche variazione di rilievo rispetto al passato nel pianificare il nuovo ciclo triennale: l'insegnamento della Fisica (termodinamica, elettromagnetismo, ottica) viene postposto al secondo e al terzo anno, quando gli studenti dovrebbero essere già in possesso delle nozioni matematiche adeguate; l'insegnamento della Meccanica propriamente detta viene poi affidato ai corsi di Fisica matematica anch'essi nel secondo e nel terzo anno. La

maggior novità è forse proprio l'istituzione di questo corso di Fisica matematica I. Esso ha uno scopo dichiaratamente ambizioso: familiarizzare gli studenti con aspetti anche abbastanza moderni della teoria dei sistemi dinamici facendo quasi del tutto a meno dell'apparato matematico di analisi, algebra e geometria che viene sviluppato nei corsi paralleli omonimi. L'uso della simulazione numerica al calcolatore, sempre più importante nello studio dei sistemi dinamici, costituisce parte integrante di questo processo di familiarizzazione. Oltre ad abituare fin da subito gli studenti a mettere le mani sul calcolo scientifico, si spera che la presentazione di questi argomenti possa contribuire a due ulteriori processi formativi di sicuro valore: da una parte, vedere nascere in modo quasi spontaneo concetti matematici profondi e sottili e vederli all'opera nel concreto; dall'altra, abituarsi fin da subito, anche se in casi semplicissimi, a lavorare con la matematica per analizzare quantitativamente le scienze della natura.

Le esercitazioni di laboratorio costituiscono, ripeto, parte insostituibile di questo corso. Esse consistono in esercizi di carattere numerico sulle sue parti fondamentali, e sono contenuti in un codice elaborato tramite il software "Mathematica" dal Dr. Mirko Degli Esposti. Il cosiddetto "Notebook" è a disposizione degli studenti nel server del Laboratorio didattico del Dipartimento di Matematica.

Ringrazio calorosamente Pierluigi Contucci, Mirko Degli Esposti, ma soprattutto Emanuela Caliceti e Michele Benzi, per avere letto queste note, correggendo numerosissimi errori e dando consigli preziosi per il miglioramento dell'esposizione.

Infine, sarò grato a chiunque leggerà queste note per segnalazioni di errori di qualsiasi natura; in particolare sarò grato agli studenti se vorranno intraprendere uno sforzo in questa direzione, come una di loro ha già cominciato a fare.

Bologna, 21.12.2001

Sandro Graffi

Indice

1	Esempi introduttivi	3
1	La formula dell'interesse composto	3
2	Ricorrenze a due termini ed incremento esponenziale	9
3	Limiti all'incremento esponenziale. La mappa logistica	14
4	Ricorrenze a tre termini: numeri di Fibonacci	16
2	Successioni, serie, frazioni continue	27
1	Generalità	27
2	Limite di una successione	29
3	Alcuni teoremi fondamentali sui limiti	34
4	Serie numeriche	37
5	Successioni di 0 e 1	43
6	Frazioni continue	55
3	Dinamica discreta	69
1	Generalità	69
2	Orbite, punti fissi, periodicità	77
3	Caratterizzazione dei punti fissi	80
4	Punti periodici e orbite periodiche	82
5	Punti fissi. Stabilità, instabilità, attrazione, repulsione	86
4	Studio della mappa logistica: frattali e caos	93
1	L'iperbolicità dei punti periodici	93
2	La mappa logistica al crescere del parametro	97
3	L'insieme ternario di Cantor	99

4	La mappa logistica genera un insieme di Cantor	103
5	La dinamica simbolica	107
6	Mappa logistica sull'insieme di Cantor e dinamica simbolica	113
7	Dipendenza delicata dalle condizioni iniziali. Caos	118
5	Biforcazioni e transizione al caos	125
1	Esempi di biforcazioni	125
2	Biforcazioni: nodo-sella e raddoppio di periodo	129
3	Caos tramite infiniti raddoppiamenti di periodo	131
6	Dinamica discreta bidimensionale	141
1	Qualche richiamo di algebra lineare	141
2	Autovalori e autovettori. Iterazione di matrici	147
3	Dinamica discreta ed iterazione di matrici 2×2	161
4	Dinamica discreta ed automorfismi del toro	168
5	Dinamica iperbolica sul toro e caos	173
6	Partizione di Markov e dinamica simbolica	176
7	Appendice	181
1	Dimostrazione del Teorema 6.7	181
2	Dimostrazione dei Teoremi 2.9,2.10,2.11	183
3	Dimostrazione delle proposizioni sulla derivata di Schwarz	185
4	Completamento della dimostrazione del Teorema 5.14	187

Capitolo 1

Esempi introduttivi

Cominciamo con lo studiare due esempi: la formula dell'interesse composto e i numeri di Fibonacci.

Arriveremo così al concetto di successione, poi a quello di ricorrenza a due e a tre termini e ne studieremo gli andamenti.

1 La formula dell'interesse composto

Si depositi in una banca la somma di denaro A ad un interesse annuo costante x .

Problema

Supponendo di non prelevare nulla né di effettuare alcun ulteriore deposito, quale sarà l'ammontare totale del deposito a seguito degli interessi composti maturati dopo n anni?

Soluzione

Dopo un anno l'interesse maturato è xA e quindi il saldo vale $A + xA = A(1 + x)$. Dopo due anni l'interesse sarà $A(1 + x)x$ e il saldo $A(1 + x) + A(1 + x)x = A(1 + x)^2$. Dopo tre anni l'interesse sarà $A(1 + x)^2x$ e il saldo $A(1 + x)^2 + A(1 + x)^2x = A(1 + x)^3$. Chiaramente dopo n anni l'ammontare totale sarà $A(1 + x)^n$.

Esercizio 1.1

Se l'interesse è il 5%, in quanti anni il capitale iniziale A raddoppierà? (Si usi una calcolatrice).

Soluzione. Dopo n anni il capitale diventerà $A(1 + 0.05)^n$ e sarà il doppio di quello iniziale quando $A(1 + 0.05)^n = 2A$ cioè $(1 + 0.05)^n = 2$. Prendendo il logaritmo naturale di ambo

i membri si trova $n = \frac{\ln 2}{\ln(1.05)}$. Ora $\ln 2 \sim 0.692$ e $\ln(1.05) \sim 0.0475$. (Il simbolo $A \sim B$ significa che il valore numerico di A è circa uguale a quello di B . Non abbiamo bisogno per il momento di precisare ulteriormente questo concetto). Quindi $n \sim 14.5$. Concludendo: il capitale raddoppierà in meno di quindici anni.

Questo fenomeno si chiama spesso raddoppio del capitale *nominale*, perché non si tiene conto della perdita del potere d'acquisto della moneta. Per il capitale vero, o capitale *reale*, quello in cui si tiene invece conto di questa perdita (detta *svalutazione*), il discorso è diverso. Tutti noi osserviamo che anno dopo anno con una quantità fissa di denaro, ad esempio 100.000 lire, si comprano sempre meno cose. A questo fenomeno si dà appunto il nome di *svalutazione della moneta*. Essa è dovuta all'inflazione, fenomeno economico sul quale non è nostro compito insistere. Ci limitiamo ad osservare che nella realtà per calcolare l'aumento (o la diminuzione) del capitale *reale* dovremo tenere conto anche del tasso di svalutazione della moneta.

Esercizio 1.2

Se il tasso di svalutazione si mantiene costantemente superiore di un punto percentuale al tasso di interesse, cosa succederà al capitale iniziale A dopo averlo lasciato in banca 15 anni alle condizioni precedenti?

Soluzione. Sia x il tasso di interesse; se il tasso di svalutazione si mantiene costantemente superiore di un punto percentuale al tasso di interesse, esso sarà $x + 0.01$. Pertanto, per il ragionamento precedente, dopo un anno il capitale A si sarà ridotto a $A = A(1 + x - x - 0.01) = A(1 - 0.01)$, dopo 2 anni a $A(1 - 0.01)^2$ e dopo 15 anni a $A(1 - 0.01)^{15}$. Ora $(1 - 0.01)^{15}$ vale circa 0.86. Pertanto il capitale iniziale perde più di un decimo del suo valore. Se il tasso di svalutazione si mantiene costantemente di due punti percentuali superiore al tasso di interesse, la diminuzione percentuale dopo 15 anni è $(1 - 0.02)^{15}$ che vale circa 0.68. Quindi dopo 15 anni il capitale iniziale perde circa un terzo del suo valore. Poiché non può succedere *mai* che il tasso di interesse sia superiore a quello di svalutazione (altrimenti le banche sarebbero sempre in perdita), ed anzi succede spesso che gli sia molto inferiore, questo spiega il perché la gente cerchi di investire i propri risparmi in vari modi ritenuti più redditizi invece che lasciarli in banca.

Definizione 1.1

1. Il prodotto $1 \cdot 2 \cdot 3 \cdots n$ dei primi n numeri naturali si denota $n!$ (e si legge n fattoriale).
Si ha $1! = 1$ e per convenzione si pone poi $0! = 1$.

2. Dato il numero naturale $0 \leq k \leq n$, si denota $\binom{n}{k}$ la frazione

$$\binom{n}{k} := \frac{n!}{k!(n-k)!} \quad (1.1)$$

e la si chiama coefficiente binomiale.

Ad esempio:

$$\begin{aligned} \binom{n}{0} &= \frac{n!}{0!(n)!} = 1; & \binom{n}{1} &= \frac{n!}{1!(n-1)!} = n; & \binom{n}{2} &= \frac{n!}{2!(n-2)!} = \frac{n(n-1)}{2}; \\ \binom{n}{k} &= \frac{n!}{k!(n-k)!} = \frac{n(n-1) \cdots (n-k+1)}{k!}; \\ \binom{n}{k} &= \frac{n!}{(n-k)!k!} = \frac{n(n-1) \cdots (n-k+1)}{k!} = \binom{n}{n-k} \\ \binom{n}{n} &= \frac{n!}{n!(0)!} = 1 = \binom{n}{0} \end{aligned}$$

Osservazione 1.1

1. I coefficienti binomiali sono tutti numeri interi positivi nonostante siano definiti dalla frazione (1.1). Essa è in realtà una frazione apparente. Infatti, poiché

$$\begin{aligned} \binom{n}{k} &= \frac{n!}{k!(n-k)!} = \frac{n(n-1) \cdots (n-k+1) \cdot (n-k) \cdots 1}{(n-k) \cdots 1 \cdot k!} \\ &= \frac{n(n-1) \cdots (n-k+1)}{k!} \end{aligned}$$

basta far vedere che $n(n-1) \cdots (n-k+1)$ è divisibile per $k! = 1 \cdot 2 \cdots k$. Controlliamo la validità di questa affermazione. Per $k = 1$ è ovvia. Per $k = 2$, il numeratore $n(n-1)$ è sicuramente un numero pari perché prodotto di due numeri consecutivi di cui uno certamente è pari; quindi è divisibile per $2! = 2$. Se $k = 3$, il numeratore è $n(n-1)(n-2)$, che è divisibile per 2 perché lo è il suo fattore $n(n-1)$, e lo è anche per 3 perché è il prodotto di tre numeri consecutivi: uno di questi è necessariamente divisibile per tre perché fra tre numeri consecutivi c'è sempre un multiplo di 3. Quindi

$n(n-1)(n-2)$ è divisibile per $1 \cdot 2 \cdot 3 = 3!$. Procedendo sempre allo stesso modo, si trova che il prodotto di k numeri consecutivi è divisibile per k . Concludendo, il prodotto dei k numeri consecutivi $n(n-1) \cdots (n-k+1)$ è divisibile per $1!2! \cdots k = k!$.

2. Si noti poi che i coefficienti binomiali non cambiano sostituendo k con $n-k$:

$$\binom{n}{n-k} = \frac{n!}{(n-k)!k!} = \binom{n}{k} : \quad (1.2)$$

Vale allora la formula detta del *binomio di Newton*¹:

$$(x+y)^n = x^n + \binom{n}{1} x^{n-1}y + \binom{n}{2} x^{n-2}y^2 + \dots + \binom{n}{n-1} xy^{n-1} + y^n \quad (1.3)$$

Usiamo l'abbreviazione di scrivere la somma degli $n+1$ numeri $a_0, a_1, \dots, a_n$ tramite il simbolo $\sum$ (detto *sommatoria*), cioè:

$$\sum_{k=0}^n a_k := a_0 + \dots + a_n$$

in cui l'indice in basso corrisponde al primo addendo e quello in alto all'ultimo. Allora potremo riscrivere la (1.3) nella sua forma abituale

$$(x+y)^n = \sum_{k=0}^n \binom{n}{k} x^k y^{n-k} \quad (1.4)$$

dove abbiamo tenuto conto del fatto che $\binom{n}{0} = \binom{n}{n} = 1$. Per $k=2$ ritroviamo la formula del quadrato del binomio $(x+y)^2 = x^2 + 2xy + y^2$, per $n=3$ quella del cubo del binomio $(x+y)^3 = x^3 + 3x^2y + 3xy^2 + y^3$, ecc. In particolare, dunque, per $y=1$:

$$(1+x)^n = \sum_{k=0}^n \binom{n}{k} x^k \quad (1.5)$$

Esercizio 1.3

Dimostrare la formula (1.4).

¹Isaac Newton (Woolsthorpe 1642- Londra 1727), lo scopritore delle leggi della meccanica e della gravitazione universale, Professore a Cambridge e Direttore della Zecca di Londra. Il suo contributo più importante alla matematica pura, ottenuto contemporaneamente a Gottfried Wilhelm Leibnitz (Lipsia 1646- Hannover 1716), ed in concorrenza poco amichevole con lui, è l'invenzione del calcolo infinitesimale. La formula del binomio di Newton chiarisce perché la frazione $\binom{n}{k}$ si chiama coefficiente binomiale.

Soluzione. Dimostriamo la formula per induzione. Anzitutto ricordiamo il *Principio dell'induzione matematica* che formuleremo così:

Sia $P(n)$ un'affermazione che dipende dal numero naturale $n \in \mathbb{N}$. Sia P vera per $n = 1$. Se la validità di P per $n = N$ implica la validità di $P(N + 1)$ allora $P(n)$ è vera $\forall n \in \mathbb{N}$.

Qui l'affermazione $P(n)$ è la validità della formula (1.4). Essa è ovviamente vera per $n = 1$. Assumendone la validità per $n = N$, dimostriamola per $n = N + 1$. In altre parole, ammettendo che sia

$$(x + y)^N = \sum_{k=0}^N \binom{N}{k} x^k y^{N-k}$$

dobbiamo dimostrare che si ha anche:

$$(x + y)^{N+1} = \sum_{k=0}^{N+1} \binom{N+1}{k} x^k y^{N+1-k}$$

Osserviamo anzitutto che è sufficiente dimostrare l'affermazione per $y = 1$: è sufficiente cioè dimostrare la (1.5). Infatti se quest'ultima è vera si ha, ponendo $t = \frac{x}{y}$:

$$\begin{aligned} (x + y)^n &= y^n (1 + t)^n = y^n \sum_{k=0}^n \binom{n}{k} t^k = \\ &= \sum_{k=0}^n \binom{n}{k} y^n \left(\frac{x}{y}\right)^k = \sum_{k=0}^n \binom{n}{k} x^k y^{n-k} \end{aligned}$$

Ricaviamo anzitutto alcune proprietà dei coefficienti binomiali.

(a)

$$\binom{n}{k} = \frac{k+1}{n+1} \binom{n+1}{k+1}, \quad k = 0, \dots, n$$

Infatti:

$$\binom{n+1}{k+1} = \frac{(n+1)!}{(k+1)!(n-k)!} = \frac{n+1}{k+1} \binom{n}{k}, \quad 0 \leq k \leq n$$

(b)

$$\binom{n}{k} = \frac{n+1-k}{n+1} \binom{n+1}{k}, \quad k = 0, \dots, n$$

Infatti:

$$\frac{n+1-k}{n+1} \binom{n+1}{k} = \frac{n+1-k}{n+1} \frac{(n+1)!}{k!(n+1-k)!} = \frac{n!}{(n-k)!k!} = \binom{n}{k}, \quad 0 \leq k \leq n$$

Possiamo ora fare vedere che

$$(1+x)^N = \sum_{k=0}^N \binom{N}{k} x^k \quad \text{implica} \quad (1+x)^{N+1} = \sum_{k=0}^{N+1} \binom{N+1}{k} x^k.$$

Si ha:

$$\begin{aligned} (1+x)^{N+1} &= (1+x)(1+x)^N = (1+x) \sum_{k=0}^N \binom{N}{k} x^k = \\ &= \sum_{k=0}^N \binom{N}{k} x^k + \sum_{k=0}^N \binom{N}{k} x^{k+1} = \\ &= \sum_{k=0}^N \frac{N+1-k}{N+1} \binom{N+1}{k} x^k + \sum_{k=0}^N \frac{k+1}{N+1} \binom{N+1}{k+1} x^{k+1} = \\ &= \sum_{k=0}^N \frac{N+1-k}{N+1} \binom{N+1}{k} x^k + \sum_{k=1}^{N+1} \frac{k}{N+1} \binom{N+1}{k} x^k = \\ &= \sum_{k=1}^N \frac{N+1}{N+1} \binom{N+1}{k} x^k + 1 + x^{N+1} = \sum_{k=0}^{N+1} \binom{N+1}{k} x^k \end{aligned}$$

e ciò conclude la prova.

Esercizio 1.4

Dimostrare la formula:

$$\sum_{k=0}^n \binom{n}{k} = 2^n, \quad n = 0, 1, 2, \dots$$

Esercizio 1.5

Verificare per calcolo diretto che i numeri $\binom{n}{k}$ soddisfano la relazione seguente (triangolo di Tartaglia². Il triangolo di Tartaglia è noto anche come triangolo di Pascal³).

1

1 1

1 2 1

1 3 3 1

²Una nota storica su Nicolò Tartaglia si trova nel capitolo successivo.

³Blaise Pascal (Clermont-Ferrand 1623-Parigi 1661), matematico, fisico e filosofo. Inventò la prima macchina calcolatrice meccanica per facilitare il lavoro del padre, Étienne Pascal, che gestiva l'esattoria delle tasse per la Normandia.

$$\begin{array}{cccccc}
 & 1 & 4 & 6 & 4 & 1 \\
 & & 1 & 5 & 10 & 10 & 5 & 1 \\
 & & & \dots & \dots & \dots & \dots & \dots
 \end{array}$$

I numeri $A(1+x)^n : n = 0, 1, 2, \dots$ generati dalla formula dell'interesse composto sono un esempio concreto di un concetto matematico molto generale, le *successioni numeriche*.

Definizione 1.2 Dicesi *successione (reale)* una qualsiasi applicazione $f : \mathbb{N} \cup \{0\}$ in $\mathbb{R}$:

$$n \mapsto f(n) = a_n.$$

Osservazione 1.2

1. Notazioni equivalenti per la successione $n \mapsto f(n) = a_n$ sono: a_n , $a(n)$, $\{a_n\}$, $(a_n)_{n \in \mathbb{N}}$, $\{a_n\}_{n=0}^\infty$, $a_0, a_1, \dots$, ecc.. Si comincia la numerazione da 0 per convenzione; sarebbe del tutto equivalente cominciarla da 1, e definire la successione come un'applicazione da $\mathbb{N}$ in $\mathbb{R}$. Nel seguito lo faremo spesso senza dirlo esplicitamente.
2. Il numero a_n si dice anche *termine generale* della successione.
3. Tutti i casi considerati in precedenza sono esempi particolari di successioni.
4. È utile tracciare il grafico di alcune successioni elementari nel piano cartesiano:

$$\begin{array}{lll}
 (a) & a_n = n; & (b) \quad a_n = \frac{1}{n}; \quad (c) \quad a_n = n^2; \\
 (d) & a_n = a^n, a > 1; & (e) \quad a_n = a^n, a < 1; \quad (f) \quad a_n = c, \forall n \in \mathbb{N}
 \end{array}$$

Le successioni formeranno l'oggetto di gran lunga principale di questo corso. Cominciamo coll'esaminarne alcune classi molto significative.

2 Ricorrenze a due termini ed incremento esponenziale

La formula dell'interesse composto è soluzione particolare del seguente problema generale:

Trovare, per $n = 1, 2, \dots$ la successione dei numeri a_n definita induttivamente nel modo seguente:

$$a_0 = \alpha; \quad a_{n+1} = \lambda a_n \tag{2.6}$$

dove λ e α sono fissati numeri reali.

In generale, una successione $a_0, a_1, a_2, \dots, a_n, \dots$ di numeri reali si dice *definita induttivamente* o *per induzione* quando

1. Viene assegnato a_0 ;
2. Si definisce a_n in termini di $a_{n-1}, a_{n-2}, \dots$, e a_0 .

La (2.6) è un esempio di successione definita induttivamente. La legge di induzione è che ogni termine è proporzionale al precedente con la medesima costante di proporzionalità. La relazione (2.6) si dice anche, per ragioni evidenti, *relazione di ricorrenza a due termini*, e il numero α *condizione iniziale* se si pone $a_0 = \alpha$. Induttivamente, si ottiene subito $a_1 = \lambda\alpha, a_2 = \lambda^2\alpha, \dots$ e in generale

$$a_n = \lambda^n \alpha, \quad n = 0, 1, 2, \dots \quad (2.7)$$

La successione a_n così ottenuta si chiama *progressione geometrica di ragione λ* .

Un altro esempio è rappresentato dalla legge di induzione seguente:

$$a_0 = \alpha; \quad a_n = a_0 + n\beta$$

Si ottiene immediatamente:

$$a_n = \alpha + n\beta. \quad (2.8)$$

Questa successione si chiama *progressione aritmetica di ragione β*

È immediato osservare quanto segue:

$$\begin{cases} \frac{|a_n|}{|a_{n-1}|} = |\lambda| > 1 & \text{se } |\lambda| > 1 \\ \frac{|a_n|}{|a_{n-1}|} = |\lambda| < 1 & \text{se } |\lambda| < 1 \end{cases} \quad (2.9)$$

mentre

$$\begin{cases} a_n = \alpha \quad \forall n \text{ se } \lambda = 1 \\ a_n = (-1)^n \alpha \quad \forall n \text{ se } \lambda = -1 \end{cases} \implies |a_n| = |\alpha| \text{ se } |\lambda| = 1 \quad (2.10)$$

In altre parole, i numeri a_n crescono o decrescono (in valore assoluto) della medesima quantità ad ogni passo. Questa proprietà si esprime dicendo che il *tasso di crescita*⁴ al passo n della

⁴In inglese: rate of increase.

successione numerica a_n , definito da $\tau(n) := \frac{|a_{n+1}|}{|a_n|}$ è costante. Più precisamente, si parlerà di tasso di crescita se $\tau(n) > 1$, e di tasso di *decrescita* se $\tau(n) < 1$. Equivalentemente, si dice che la successione a_n soddisfa la *legge dell'incremento esponenziale* (se $|\lambda| > 1$) o del *decremento esponenziale* (se $|\lambda| < 1$).

Dunque le locuzioni di incremento (o decremento) esponenziale o *geometrico* sono sinonimi.

Osservazione 2.3

1. È bene notare che la soluzione (2.7) della legge di ricorrenza a due termini definita induttivamente dalla (2.6) costituisce l'esempio più semplice di *evoluzione deterministica*: data la legge (2.6), la conoscenza del *dato iniziale* α determina in modo univoco tutti gli elementi a_n .
2. L'incremento esponenziale è una crescita molto "rapida". Facciamo qualche esempio numerico di confronto con crescite di altro tipo, ad esempio con la successione a_n definita dalla legge $a_n := n^p, n = 0, 1, \dots$ dove $p > 0$. Per definizione, si tratta della legge di incremento di potenza di esponente p . $a_n = n^{-p}$ definisce la legge di decremento di potenza di esponente p . Per confrontare le due crescite confrontiamo il loro tasso di crescita, cioè il rapporto fra ogni termine e quello che lo precede. Sappiamo che il tasso di crescita è costante per la legge esponenziale. Calcoliamolo per la legge di potenza. Si ha

$$\frac{|a_n|}{|a_{n-1}|} = \frac{(n+1)^p}{n^p} = \frac{n^p(1+1/n)^p}{n^p} = \left(1 + \frac{1}{n}\right)^p$$

D'altra parte $(1 + 1/n)^p$ decresce all'aumentare di n : infatti

$$\left(1 + \frac{1}{(n+1)}\right)^p < \left(1 + \frac{1}{n}\right)^p \quad \text{dato che} \quad \frac{1}{(n+1)} < \frac{1}{n}.$$

Dunque il tasso di crescita *diminuisce* con n , e al crescere di n diventerà prima o poi più piccolo di qualsiasi costante $\lambda > 1$ che è il tasso di crescita esponenziale crescente. Pertanto la crescita di potenza è meno rapida di qualunque crescita esponenziale. Facendo il ragionamento analogo per il decremento si conclude che il decremento esponenziale è più veloce di quello di qualsiasi legge di potenza reciproca $\frac{1}{n^p}, p > 0$.

3. Supponendo per comodità $a_n > 0$ fissiamo un sistema di assi cartesiani in cui nelle ascisse riportiamo n e nelle ordinate $\ln n$. Si dice in tal caso che si sta usando una *scala logaritmica*. (Con $\ln n$, o talvolta $\log n$ indichiamo il *logaritmo naturale* del numero n , cioè il logaritmo in base e , dove e è il numero di Nepero. Comunque la scala logaritmica può essere definita mettendo in ordinata il logaritmo di n in qualsiasi base). Riportiamo ora i punti λ^n in tale scala. Se in ascissa riportiamo n , in ordinata riporteremo $n \ln \lambda$. I punti si ottengono quindi aggiungendo sempre la medesima quantità cioè $\ln \lambda$, al punto precedente. Ricordiamo ora che si dice *progressione aritmetica di ragione α* ogni successione definita induttivamente da $a_1 = \beta, a_{n+1} = a_n + \alpha$, dove $\beta \in \mathbb{R}$ è arbitrario. In tal modo si ha $a_2 = \alpha + \beta, a_3 = 2\alpha + \beta, \dots, a_{n+1} = n\alpha + \beta, \dots$. I punti della successione $a_n = \lambda^n$ si dispongono quindi in una progressione aritmetica di ragione $\ln \lambda$. In altre parole, il logaritmo trasforma progressioni geometriche in progressioni aritmetiche. La ragione sarà positiva se $\ln \lambda > 0$ cioè $\lambda > 1$, e dunque se la ragione della progressione geometrica è maggiore di 1; sarà negativa se $\ln \lambda < 0$ cioè $\lambda < 1$, e dunque se la ragione della progressione geometrica è minore di 1. Essi stanno quindi sulla retta di equazione $y = \ln \lambda x$ di coefficiente angolare $\ln \lambda$, e pertanto crescono *linearmente*. Viceversa, se in scala logaritmica vediamo che i dati si dispongono in progressione aritmetica, concluderemo che la successione cresce o decresce esponenzialmente, cioè come λ^n con $\lambda > 1$ o $0 < \lambda < 1$. In altre parole la scala logaritmica trasforma crescite esponenziali in crescite lineari, molto più facili da riconoscere a prima vista in un grafico.

Esercizio 2.6

1. Si considerino le successioni $a_n = n! := 1 \cdot 2 \cdot 3 \cdots n$; $b_n = \lambda^n$. Si dimostri applicando il ragionamento precedente che il tasso di crescita di a_n è maggiore di quello di b_n (*il fattoriale cresce più velocemente di qualsiasi esponenziale*). (Da un certo n in poi).
2. Si considerino le successioni $a_n = n! := 1 \cdot 2 \cdot 3 \cdots n$; $b_n = n^n$. Applicando il ragionamento precedente dimostrare che il tasso di crescita di a_n è minore di quello di b_n . (*n^n cresce più velocemente del fattoriale*).
3. Dimostrare che l'esponenziale trasforma progressioni aritmetiche in progressioni geo-

metriche. In altre parole, dimostrare che se $a_{n+1} = a_n + \beta, n = 0, 1, \dots$, allora

$$\frac{e^{a_{n+1}}}{e^{a_n}} = \alpha$$

e calcolare la ragione α .

Esempio 2.1

La legge dell'interesse composto è un esempio di incremento esponenziale ($\lambda = 1 + x$). La legge della perdita di valore della moneta a tasso di svalutazione costante è un esempio di decremento esponenziale ($\lambda = 1 - x$). Esempi significativi di incrementi come n^n li incontreremo nel capitolo successivo.

Esercizio 2.7

Assumendo che una generazione duri in media 25 anni, trovare il numero massimo degli antenati di ciascuno di noi nell'anno mille.

Soluzione. n generazioni fa ciascuno di noi aveva 2^n antenati (2 genitori, 4 nonni, 8 bisnonni, ecc.); costoro non erano però a priori tutti distinti, cosicché 2^n è a priori solo il loro numero massimo. Dall'anno mille a oggi sono trascorse 40 generazioni, per cui ciascuno di noi avrebbe potuto avere al più 2^{40} antenati.

Si noti però che $2^{40} = (2^{10})^4 = (1024)^4 > 1000^4 = 10^{12} = 1000$ miliardi. La popolazione mondiale nell'anno mille era sicuramente meno di 500 milioni di persone. Dunque il risultato è palesemente assurdo. Già nel 1450, ad esempio, cioè 22 generazioni fa, ciascuno di noi avrebbe dovuto avere $2^{22} > 1$ miliardo di antenati. Basterebbe dunque risalire al 1450, quando il numero totale degli abitanti della terra era ben lontano dal raggiungere il miliardo (raggiunto solo verso il 1900), per concludere che ciascuno di noi è parente di ciascun giapponese, ciascun indio della Terra del Fuoco e ciascun aborigeno australiano. La soluzione di questo apparente paradosso è evidente: i nostri antenati non sono tutti distinti. Anzi, queste stime rozze mostrano chiaramente come la molteplicità dei nostri antenati sia alta, ed in ultima analisi quanto la composizione etnica della popolazione mondiale sia stabile.

3 Limiti all'incremento esponenziale. La mappa logistica

Il tasso di incremento costante ad ogni passo che produce l'incremento esponenziale è un'approssimazione molto rozza per questioni tipo l'aumento della popolazione di una specie e la diffusione delle epidemie. Sia p_n la popolazione dopo n generazioni, oppure il numero di persone contagiate da una malattia qualsiasi dopo n anni. Domanda:

Come varierà la popolazione alla generazione successiva, o come aumenterà il contagio l'anno seguente?

La legge più semplice che si possa immaginare è che l'aumento sia direttamente proporzionale alla popolazione presente (nel caso del contagio, a quella dei malati). Dunque si suppone $p_{n+1} = kp_n$; $k > 0$ da cui segue come sappiamo la legge esponenziale $p_n = k^n p_0$. Ora qui si possono avere due casi non banali: $k > 1$ o $k < 1$ (per $k = 1$ si ha banalmente $p_n = p_0$ per ogni n). Nel primo caso la popolazione cresce esponenzialmente, e il contagio anche; nel secondo la popolazione decresce esponenzialmente, e il contagio tende ad esaurirsi. Non ci sono altre possibilità. Visto che le cose non vanno così (l'esempio più recente è costituito dal contagio dell'AIDS: all'inizio si temeva l'incremento esponenziale, poi si è visto che il tasso di crescita è fortemente diminuito, salvo che in alcuni paesi africani dove le contromisure sono praticamente inesistenti), è chiaro che un simile modello è troppo rozzo. Oggi si sa che può descrivere, nel caso delle epidemie, al più la loro fase iniziale.

Un modello un pò meno rozzo, che si applica alla crescita delle popolazioni, incorpora un'assunzione in più: quella che la crescita non possa superare un certo valore, detto *valore limite*. L'esistenza di questo valore può essere dovuto alle cause più diverse, quali la disponibilità di cibo, di spazio, l'equilibrio con altre specie, ecc., sulle quali non è il caso di entrare qui. Sia L questo valore limite. Passando eventualmente alle percentuali di popolazione di una specie rispetto a tutte o parte delle altre, potremo sempre assumere $L = 1$. Dire che L è un valore limite significa che la popolazione dovrà decrescere dopo avere raggiunto questo valore. D'altra parte, ci si aspetta crescita finché la popolazione non ha raggiunto il valore L . Assumendo senza ledere la generalità $L = 1$ il modello più semplice che incorpora queste due richieste è il seguente:

$$p_{n+1} = kp_n(1 - p_n) \tag{3.11}$$

dove k è una costante numerica che può assumere valori diversi a seconda delle circostanze. La "legge" (3.11) viene chiamata *mappa logistica*. Si tratta ancora di una ricorrenza a due termini, e quindi l'evoluzione è completamente determinata una volta fissato il dato iniziale p_0 . Si noti tuttavia che essa è quadratica (cioè è definita da una funzione polinomiale di secondo grado), e quindi notevolmente più complicata di quella lineare che si risolve subito tramite elevamento a potenza come abbiamo appena visto.

La ricorrenza dà invece origine ad una composizione n -esima della medesima funzione. Sia infatti $f(p) := kp(1-p)$. Allora la (3.11) può essere considerata come l'iterazione n -esima di f , cioè la composizione n -esima di f con sè stessa. Infatti la ricorrenza

$$p_1 = f(p_0) = kp_0(1-p_0); p_2 = f(p_1) = kp_1(1-p_1); p_3 = f(p_2) = kp_2(1-p_2), \dots$$

può essere riscritta così (si faccia attenzione a non cadere nell'errore frequente di confondere la composizione di una funzione con sè stessa n volte con con l'elevamento della funzione medesima alla potenza n):

$$p_1 = f(p_0); p_2 = (f \circ f)(p_0); p_3 = (f \circ f \circ f)(p_0); \dots, p_n = (f \circ f \circ f \dots \circ f)(p_0)$$

A questo proposito si ricordi che se il codominio J della funzione $f(x)$ definita su un intervallo $I_1 \subset \mathbb{R}$ è contenuto nell'intervallo $I_2 \subset \mathbb{R}$ sul quale supponiamo definita la funzione $g(x)$, allora la funzione composta $g \circ f$ è definita così: $(g \circ f)(x) = g(f(x))$, $\forall x \in I_1$. Ad esempio:

$$\begin{aligned} f &: \mathbb{R} \rightarrow [-1, 1], x \mapsto f(x) = \cos x; I_1 = \mathbb{R}, J = [-1, 1]; \\ g &: \mathbb{R} \rightarrow \mathbb{R}_+, x \mapsto \frac{1}{1+x^2}; I_2 = \mathbb{R}; \\ g \circ f &: \mathbb{R} \rightarrow \mathbb{R}_+, x \mapsto (g \circ f)(x) = \frac{1}{1+\cos^2 x} \end{aligned}$$

(Un altro errore frequente da evitare è ritenere che l'operazione di composizione fra due funzioni sia commutativa: $g \circ f \neq f \circ g$!. Nell'esempio precedente infatti $(f \circ g)(x) = \cos\left(\frac{1}{1+x^2}\right)$. Addirittura $f \circ g$ può essere definita mentre $g \circ f$ no. Esempio: $f(x) = \sqrt{x}$, definita per $x > 0$, e $g(x) = -3$ definita $\forall x \in \mathbb{R}$. Allora $(g \circ f)(x) = -3$ mentre $(f \circ g)$ non esiste).

Vedremo nel Cap.4 che le soluzioni di questa ricorrenza dipendono in maniera estremamente delicata dal dato iniziale p_0 , nel senso che dati iniziali vicini a piacere possono dare origine ad evoluzioni radicalmente diverse. Ciò fa sì che l'insieme delle sue soluzioni sia assai ricco di strutture interessanti.

4 Ricorrenze a tre termini: numeri di Fibonacci

Abbiamo visto degli esempi di ricorrenze a due termini, con incremento costante o no. Esse sono leggi deterministiche, nel senso che permettono di generare la soluzione conoscendo solo il termine iniziale e danno origine alla legge dell'incremento o del decremento esponenziale quando l'incremento o il decremento sono costanti.

Trattiamo ora le ricorrenze a tre termini, sempre con incrementi costanti, che sono ancora leggi deterministiche. Cominciamo con un celebre esempio (contenuto nel *Liber abaci* di Leonardo Fibonacci⁵).

Esempio 4.2

Un uomo colloca una coppia di conigli in un luogo isolato da ogni lato da un muro. Quante coppie di conigli saranno generate a partire da questa nell'arco di un anno supponendo che ogni nuova coppia diventi fertile dopo due mesi e che ogni coppia fertile ne generi una nuova ogni mese?

Impostazione della soluzione. Sia a_n il numero di coppie di conigli al mese n . Al mese $n + 2$ il numero a_{n+2} delle coppie di conigli sarà il numero a_{n+1} delle coppie presenti al mese precedente aumentato dal numero delle coppie nuove nate, uguale al numero delle coppie che hanno generato una nuova coppia. Per essere in grado di generare una coppia deve essere fertile; ogni coppia impiega due mesi a raggiungere la fertilità; dunque il numero delle coppie fertili è a_n . Pertanto il numero a_n soddisfa la relazione $a_{n+2} = a_{n+1} + a_n$.

Il problema ammette quindi la seguente formulazione generale:

Problema. (Numeri di Fibonacci). Trovare la successione di numeri naturali $a_0, a_1, a_2, a_3 \dots$ in cui ogni elemento è la somma dei due precedenti. I primi due termini della successione sono 0 e 1.

Soluzione. Denotiamo a_n il generico numero naturale cercato. Allora la condizione che a_n

⁵Leonardo Fibonacci, o Leonardo Pisano, autore del *Liber abaci*, dal quale sono tratti gli esempi seguenti, fu un matematico attivo a Pisa nella prima metà del XIII secolo, vissuto fra il 1170 e il 1250. Fu il primo a portare in Occidente la matematica araba, all'epoca la più sviluppata perché anche non aveva mai perso il contatto con la tradizione ellenistica, che aveva appreso in vari anni di soggiorno ad Algeri praticando il commercio. È documentata anche la sua presenza alla corte di Federico II.

deve essere la somma dei due precedenti si scrive così:

$$a_{n+2} = a_{n+1} + a_n \quad (4.12)$$

Per ipotesi, $a_0 = 0$, $a_1 = 1$. Calcoliamo a titolo di esempio alcuni termini della successione: $a_2 = 1, a_3 = 2, a_4 = 3, a_5 = 5, a_6 = 8, a_7 = 13, a_8 = 21, a_9 = 34, a_{10} = 55, a_{11} = 89, \dots$. Che tipo di crescita ha questa successione? Se calcoliamo il rapporto fra ogni termine e quello precedente otteniamo:

$$\begin{aligned} \frac{a_2}{a_1} &= 1, \quad \frac{a_3}{a_2} = 2, \quad \frac{a_4}{a_3} = \frac{3}{2} = 1.5, \quad \frac{a_5}{a_4} = \frac{5}{3} = 1.6666\dots, \quad \frac{a_6}{a_5} = \frac{8}{5} = 1.6, \quad \frac{a_7}{a_6} = \frac{13}{8} = 1.625\dots, \\ \frac{a_8}{a_7} &= \frac{21}{13} = 1.615\dots, \quad \frac{a_9}{a_8} = \frac{34}{21} = 1.619\dots, \quad \frac{a_{10}}{a_9} = \frac{55}{34} = 1.618\dots, \quad \frac{a_{11}}{a_{10}} = \frac{89}{55} = 1.618\dots \end{aligned}$$

Questi rapporti tendono ad una costante maggiore di 1; l'indicazione che se ne trae è un incremento esponenziale. Cerchiamo allora la soluzione a_n ancora sotto forma di legge di incremento esponenziale. Poniamo a tal fine $a_n = \lambda^n$ e cerchiamo di determinare λ . Sostituendo $a_n = \lambda^n$, $a_{n+1} = \lambda^{n+1}$, $a_{n+2} = \lambda^{n+2}$ nella (4.12) troviamo subito la condizione di esistenza del numero λ voluto:

$$\lambda^{n+2} = \lambda^{n+1} + \lambda^n \implies \lambda^2 = \lambda + 1 \quad (4.13)$$

L'equazione di secondo grado $\lambda^2 - \lambda - 1 = 0$ ha due radici reali λ_1 e λ_2 , date da

$$\lambda_{1,2} = \frac{1 \pm \sqrt{5}}{2} \quad (4.14)$$

Pertanto ambedue le successioni λ_1^n e λ_2^n , $n = 0, 1, \dots$ soddisfano la ricorrenza (4.12). Siano α_1 e α_2 numeri reali arbitrari. Allora si vede subito che anche tutte le successioni b_n della forma $b_n := \alpha_1 \lambda_1^n + \alpha_2 \lambda_2^n$ soddisfano la ricorrenza. Infatti:

$$\begin{aligned} b_{n+2} &= \alpha_1 \lambda_1^{n+2} + \alpha_2 \lambda_2^{n+2} = \alpha_1 (\lambda_1^{n+1} + \lambda_1^n) + \alpha_2 (\lambda_2^{n+1} + \lambda_2^n) \\ &= \alpha_1 \lambda_1^{n+1} + \alpha_2 \lambda_2^{n+1} + \alpha_1 \lambda_1^n + \alpha_2 \lambda_2^n = b_{n+1} + b_n \end{aligned}$$

Tutte le successioni della forma $\alpha_1 a_n + \alpha_2 b_n$ con α_1 e α_2 numeri reali arbitrari si dicono *combinazioni lineari* delle successioni a_n e b_n . Allora possiamo affermare che tutte le combinazioni lineari delle soluzioni λ_1^n e λ_2^n sono ancora soluzioni. Questa è una proprietà fondamentale delle equazioni lineari come la (4.12), sulla quale torneremo. Equazioni del tipo (4.12) si

dicono lineari perché sono espressioni di primo grado nelle a_n . Ad esempio, l'equazione $a_{n+1} = \lambda a_n^2$ non è lineare, e nemmeno lo è l'equazione $a_{n+2} = a_n a_{n+1}$. (Poiché un'equazione di primo grado nelle due variabili (x, y) del tipo $Ax + By + C = 0$ definisce una retta nel piano cartesiano, equazioni di tal fatta si dicono lineari). Ora bisogna determinare quali fra le infinite successioni b_n soddisfano la ricorrenza di Fibonacci: a tale scopo imponiamo ad ogni successione b_n di soddisfare le condizioni iniziali richieste, e cioè $b_0 = 0, b_1 = 1$. Poiché $b_0 = \alpha_1 + \alpha_2$ e $b_1 = \alpha_1 \lambda_1 + \alpha_2 \lambda_2$ si ottengono le due condizioni

$$\begin{cases} \alpha_1 + \alpha_2 = 0 \\ \alpha_1 \lambda_1 + \alpha_2 \lambda_2 = 1 \end{cases}$$

Si tratta di un sistema lineare di due equazioni nelle due incognite α_1 e α_2 . La prima equazione implica $\alpha_2 = -\alpha_1$, e la seconda $\alpha_1 = \frac{1}{\lambda_1 - \lambda_2} = \frac{1}{\sqrt{5}}$. Questa coppia di valori di α_1 e α_2 è l'unica soluzione del sistema. Pertanto la successione

$$\begin{aligned} a_n &:= \frac{1}{\lambda_1 - \lambda_2} (\lambda_1^n - \lambda_2^n) = \frac{1}{2^n \sqrt{5}} [(\sqrt{5} + 1)^n - (1 - \sqrt{5})^n] \\ &= \frac{1}{\sqrt{5}} \left[\left(\frac{\sqrt{5} + 1}{2} \right)^n - \left(\frac{1 - \sqrt{5}}{2} \right)^n \right] \end{aligned} \quad (4.15)$$

soddisfa la ricorrenza $a_{n+2} = a_{n+1} + a_n$ con $a_0 = 0, a_1 = 1$ per costruzione. Dunque a_n è l' n -esimo numero di Fibonacci. (Si noti che i numeri a_n sono comunque interi nonostante la comparsa degli irrazionali $\sqrt{5}$: le potenze pari di $\sqrt{5}$ nella differenza si cancellano, e le potenze dispari danno un numero intero dopo essere state divise per $1/\sqrt{5}$ a fattor comune). È chiaro anche che la successione a_n che risolve il problema è unica: se ce ne fosse un'altra, denotata c_n , i primi due termini devono comunque essere $c_0 = 0$ e $c_1 = 1$; poiché la relazione di ricorrenza è la stessa, $c_{n+2} = c_{n+1} + c_n$ e quindi $c_n = a_n$ per ogni n .

Calcoliamo infine il numero totale delle coppie di conigli generate dopo un anno a partire dall'unica coppia; Invece di fare un calcolo esatto, facciamo una rapida stima numerica basata sulla (4.15), nella quale prenderemo $n = 12$. Anzitutto notiamo che $|\frac{1 - \sqrt{5}}{2}| \sim 0.62$, mentre $|\frac{1 + \sqrt{5}}{2}| \sim 1.62$. (Si osservi che questo numero coincide, a meno della terza cifra decimale, con la legge di incremento che avevamo ottenuto esaminando la crescita dei primo 12 termini). Pertanto, dato che $n = 12$, potremo senz'altro trascurare il secondo esponenziale rispetto al primo. Dunque:

$$a_{12} \sim \frac{(1.62)^{12}}{2.24} > \frac{(1.6)^{12}}{2.24} \sim 140$$

Concludiamo dunque che in un anno saranno generate almeno 140 coppie di conigli.

Osserviamo che dopo due anni ci saranno almeno $\frac{(1.6)^{24}}{2.24^2} = 140 \times 140 = 19600$ coppie, dopo 3 anni $19600 \times 140 = 2.744.000$ coppie, dopo 4 anni $2.744.000 \times 140 = 384.160.000$ coppie, dopo 5 anni più di 53 miliardi. Non si deve pensare che questo modello sia del tutto irrealistico: un fenomeno del genere si verificò in Australia attorno al 1840 quando vi fu importata una coppia di conigli, animali fino ad allora sconosciuti in quel continente. Essi trovarono un *habitat* talmente favorevole (spazio e cibo a volontà, assenza di predatori, ecc.) che presero a riprodursi con legge esponenziale devastando in pochi anni tutte le colture. Fu necessario importare dei predatori come le volpi per porre un limite alla proliferazione dei conigli. Vedremo in seguito il modello più semplice per descrivere il fenomeno della limitazione alla crescita esponenziale.

Osservazione 4.4

L'equazione di secondo grado $\lambda^2 - \lambda - 1 = 0$ è ben nota fin dall'antichità per un motivo molto diverso. La si ottiene infatti uguagliando il prodotto dei medi con quello degli estremi nella *divina proporzione*

$$x : a = a : (x - a) \implies x^2 - ax - a^2 = 0 \implies x_{1,2} = a \frac{1 \pm \sqrt{5}}{2}$$

che in parole si enuncia così: il tutto x sta alla parte a come la parte a sta al tutto meno la parte $x - a$. La soluzione $x_1 = a \frac{1 + \sqrt{5}}{2}$ definisce $a = \frac{2}{\sqrt{5} + 1} x_1 = \frac{\sqrt{5} - 1}{2} x_1$ come la *sezione aurea* del segmento x_1 . Equivalentemente, il numero irrazionale $\frac{\sqrt{5} - 1}{2}$ si chiama talvolta *media aurea*, o *numero aureo*. La divina proporzione si chiama così perchè l'altezza del rettangolo che costituisce la pianta di ogni tempio greco è la sezione aurea della base. L'irrazionalità della sezione aurea del segmento fu un'altra delle sensazionali scoperte dei pitagorici, al pari di quella dell'irrazionalità del rapporto fra diagonale e lato di un quadrato.

Osservazione 4.5

Con il medesimo metodo seguito per trovare i numeri di Fibonacci si risolvono *tutte* le relazioni di ricorrenza a tre termini ad incrementi costanti (o, equivalentemente, a coefficienti costanti) purché siano note le condizioni iniziali, cioè i primi due numeri della successione che si cerca.

Si tenga però ben presente che questi tipi di ricorrenza sono comunque casi molto particolari delle ricorrenze: in altre parole non c'è da sperare che questo metodo funzioni per ricorrenze diverse da quelle a tre o più termini (che vedremo dopo) a coefficienti costanti.

Una *relazione di ricorrenza a tre termini a coefficienti costanti* fra numeri reali $A_n, n = 0, 1, \dots$ è definita in generale nel modo seguente

$$A_{n+2} = aA_{n+1} + bA_n \quad (4.16)$$

dove a e b sono assegnati numeri reali, per semplicità positivi.

Problema

Determinare il termine generico A_n sapendo che $A_0 = p, A_1 = q$, dove p, q sono assegnati numeri reali.

Soluzione

Si procede come nel caso dei numeri di Fibonacci, cercando la soluzione sotto la forma $A_n = \lambda^n$. Sostituendo nella (4.16) si trova l'equazione di secondo grado (detta *equazione caratteristica* della ricorrenza)

$$\lambda^2 - a\lambda - b = 0 \implies \lambda_{1,2} = \frac{a \pm \sqrt{a^2 + 4b}}{2}$$

Poiché $b > 0$ ambedue le radici sono reali, e ovviamente distinte: $\lambda_1 \neq \lambda_2$. Si considerino allora tutte le successioni b_n della forma generale $b_n := \alpha_1 \lambda_1^n + \alpha_2 \lambda_2^n$. Ripetendo senza alcuna variazione il ragionamento fatto per la successione di Fibonacci, concludiamo che anche tutte queste soddisfano la ricorrenza (si noti che la ricorrenza (4.16) è lineare per ogni scelta della coppia (a, b)). Determiniamo ancora α_1 e α_2 richiedendo alla successione b_n di soddisfare le condizioni $b_0 = p, b_1 = q$. Imponendo queste uguaglianze ai termini con $n = 0$ e $n = 1$ della forma generale troviamo il sistema lineare di due equazioni nelle due incognite α_1 e α_2 :

$$\begin{cases} \alpha_1 + \alpha_2 = p \\ \alpha_1 \lambda_1 + \alpha_2 \lambda_2 = q \end{cases}$$

la cui unica soluzione è la coppia

$$\begin{aligned} \alpha_1 &= \frac{q - p\lambda_2}{\lambda_1 - \lambda_2} = \frac{q - p(a - \sqrt{a^2 + 4b})/2}{\sqrt{a^2 + 4b}}; \\ \alpha_2 &= p - \alpha_1 = p - \frac{q - p(a - \sqrt{a^2 + 4b})/2}{\sqrt{a^2 + 4b}} \end{aligned}$$

Una volta determinati α_1 e α_2 la soluzione per il termine generico A_n sarà:

$$A_n = \frac{q - p(a - \sqrt{a^2 + 4b})/2}{\sqrt{a^2 + 4b}} \left[\frac{a + \sqrt{a^2 + 4b}}{2} \right]^n + \left(p - \frac{q - p(a - \sqrt{a^2 + 4b})/2}{\sqrt{a^2 + 4b}} \right) \left[\frac{a - \sqrt{a^2 + 4b}}{2} \right]^n \quad (4.17)$$

Osservazione 4.6

1. Osserviamo esplicitamente che la soluzione appena scritta è unica una volta assegnati i due dati iniziali p, q . Il ragionamento fatto per la successione di Fibonacci vale in generale: se ce ne fosse un'altra, i primi due termini devono coincidere con p, q e quindi per ricorrenza anche tutti gli altri devono coincidere. La formula precedente per A_n mostra anche la dipendenza esplicita dai dati iniziali p e q .
2. Il primo esponenziale della (4.17) scompare se $\alpha_1 = 0$ cioè se $q = p(a - \sqrt{a^2 + 4b})/2$ ovvero $q = \lambda_2 p$. Il secondo invece scompare se $\alpha_2 = 0$ cioè se $q = \lambda_1 p$. Sappiamo dalla geometria analitica che le due equazioni $q = \lambda_1 p$ e $q = \lambda_2 p$ rappresentano due rette passanti per l'origine nel piano cartesiano $\mathbb{R}^2$ in cui denotiamo le ascisse p e le coordinate q , rispettivamente. I coefficienti angolari di queste due rette sono λ_1 e λ_2 , rispettivamente. Poiché $\lambda_1 \cdot \lambda_2 = -b$, se $b = 1$ queste due rette sono anche ortogonali, perché, ricordiamo, condizione necessaria e sufficiente affinché due rette siano ortogonali è che il prodotto dei loro coefficienti angolari valga -1 .

Osservazione 4.7

Naturalmente la legge dell'incremento esponenziale determinata dalla relazione di ricorrenza a due termini che abbiamo visto in precedenza è un caso particolare di questa trattazione. Data infatti la relazione di ricorrenza $A_{n+1} = aA_n$, con $A_0 = p$, ponendo $A_n = \lambda^n$ e sostituendo si ha subito $\lambda = a$ e quindi $A_n = a^n p$.

Si noti però che ciò richiede che la relazione di ricorrenza abbia luogo fra due termini successivi. Saltandone uno l'incremento (o il decremento) esponenziale può sparire e può apparire un comportamento oscillatorio. Si consideri ad esempio la ricorrenza $a_{n+2} = -a_n$. Sia $a_0 = p, a_1 = q$. Allora si ha subito $a_2 = -p, a_3 = -q, a_4 = p, a_5 = q \dots$ e in generale

$$a_{2n} = (-1)^n p, \quad a_{2n+1} = (-1)^n q \quad (4.18)$$

Per chi conosce i numeri complessi. Possiamo considerare la ricorrenza $a_{n+2} = -a_n$ come un caso particolare della (4.16) con $a = 0$, $b = -1$. Il procedimento di soluzione precedente si applica direttamente anche in questo caso. L'equazione caratteristica, ottenuta inserendo $a_n = \lambda^n$, $a_{n+2} = \lambda^{n+2}$ nella $a_{n+2} = -a_n$ è $\lambda^2 + 1 = 0$ con le radici $\lambda_{1,2} = \pm i$. Allora avremo la soluzione generale $a_n = \alpha_1 i^n + \alpha_2 (-i)^n = \alpha_1 e^{in\pi/2} + \alpha_2 e^{-in\pi/2}$. Ora $a_0 = \alpha_1 + \alpha_2$ e $a_1 = i\alpha_1 - i\alpha_2$. Quindi se imponiamo le condizioni iniziali troviamo $\alpha_1 + \alpha_2 = p$ e $i\alpha_1 - i\alpha_2 = q$ da cui

$$\alpha_1 = \frac{3p + iq}{2}, \quad \alpha_2 = \frac{-p - iq}{2}$$

e pertanto

$$a_n = \frac{1}{2} [e^{in\pi/2}(3p + iq) + e^{-in\pi/2}(-p - iq)] \Rightarrow a_{2n} = (-1)^n p, \quad a_{2n+1} = (-1)^n q.$$

Si vede subito che questa successione, lungi dal crescere o decrescere esponenzialmente, si ripete periodicamente: $a_{2n+2} = (-1)^{n+2} p = a_{2n}$, $a_{2n+3} = (-1)^{n+2} q = a_{2n+1}$.

Non si deve però pensare che basti saltare un termine nella ricorrenza per fare sparire l'incremento esponenziale. Consideriamo infatti la ricorrenza $a_{n+2} = ba_n$, $b > 0$; l'equazione caratteristica (al solito si pone $a_n = \lambda^n$, $a_{n+2} = \lambda^{n+2}$ e si sostituisce nella ricorrenza) è $\lambda^2 - b = 0$ da cui $\lambda_{1,2} = \pm\sqrt{b}$. La soluzione generale è $a_n = \alpha_1 b^{n/2} + \alpha_2 (-1)^n b^{n/2}$; quindi comunque si scelgano le condizioni iniziali, cioè per ogni scelta di α_1, α_2 , vi sarà incremento esponenziale se $b > 1$ e decremento esponenziale se $0 < b < 1$.

Il procedimento precedente si può generalizzare immediatamente alla relazione di ricorrenza a p termini, $p > 2$, definita nel modo seguente:

$$A_{n+p} = a_1 A_{n+p-1} + a_2 A_{n+p-2} + \dots + a_p A_n \quad (4.19)$$

dove $a_1, \dots, a_p$ sono assegnati numeri reali. Siano poi assegnate le condizioni iniziali, cioè i primi p termini della successione A_n : $A_0 = q_0, A_1 = q_1, \dots, A_{p-1} = q_{p-1}$. Facendo al solito la posizione $A_n = \lambda^n$ e sostituendo nella (4.19) si ottiene la corrispondente equazione caratteristica

$$\lambda^p - a_1 \lambda^{p-1} + \dots - a_{p-1} \lambda - a_p = 0 \quad (4.20)$$

Si tratta dell'equazione algebrica di grado p nella sua forma più generale.

Ricordiamo dall'algebra i seguenti importantissimi risultati:

1. Sia $Q_p(\lambda)$ un polinomio di grado p sul campo complesso $\mathbb{C}$. (Possiamo sempre assumere, senza ledere la generalità, che il coefficiente di λ^p sia 1). Allora $\lambda_1 \in \mathbb{C}$ è una *radice di molteplicità* ν_1 , $1 \leq \nu_1 \leq p$ o equivalentemente uno *zero* (di pari molteplicità) del polinomio $Q_p(\lambda)$, cioè $Q_p(\lambda_1) = 0$, se $Q_p(\lambda)$ è divisibile per $(\lambda - \lambda_1)^{\nu_1}$. Ciò significa che esiste ed è unico un polinomio $R_p(\lambda)$ di grado $p - \nu_1$ con $R_p(\lambda_1) \neq 0$ tale che $Q_p(\lambda) = (\lambda - \lambda_1)^{\nu_1} R_p(\lambda)$. Le radici di molteplicità 1 si dicono *semplici*.
2. (Teorema fondamentale dell'algebra): Ogni polinomio di grado p , a coefficienti reali o complessi, ammette esattamente p radici nel campo complesso $\mathbb{C}$. Ogni radice è contata un numero di volte pari alla sua molteplicità. Denotate $\lambda_1, \dots, \lambda_k$ le radici distinte, $1 \leq k \leq p$, con molteplicità rispettive $\nu_1, \dots, \nu_k$, si ha $\nu_1 + \dots + \nu_k = p$ e (assumendo sempre che il coefficiente di λ^n valga 1) vale la fattorizzazione

$$Q_p(\lambda) = (\lambda - \lambda_1)^{\nu_1} \cdot (\lambda - \lambda_2)^{\nu_2} \dots (\lambda - \lambda_k)^{\nu_k}$$

3. Se i coefficienti di Q_p sono tutti reali, e μ è una sua radice, $Q_p(\mu) = 0$, lo è anche la sua complessa coniugata $\bar{\mu}$ perché $Q_p(\bar{\mu}) = \overline{Q_p(\mu)}$ e $\overline{Q_p(\mu)} = 0$ se $Q_p(\mu) = 0$. In altre parole: *Le radici di un polinomio a coefficienti reali che non sono reali sono coppie di numeri complessi coniugati.*
4. Le equazioni del tipo $f(\lambda) = 0$ dove $f(\lambda)$ non è un polinomio ma una funzione con certe proprietà di regolarità che non è necessario precisare qui sono dette *trascendenti*. Esempi di equazioni trascendenti sono le equazioni trigonometriche. A differenza delle equazioni algebriche, le equazioni trascendenti possono anche non ammettere alcuna soluzione, od ammetterne infinite. Ad esempio, l'equazione esponenziale $e^\lambda = 0$ non ha soluzioni nel campo complesso. L'equazione $\sin(\lambda) = 0$ ammette le infinite soluzioni $\lambda_n = n\pi$, $n \in \mathbb{Z}$.

Dunque l'equazione $Q_p(\lambda) = 0$ ammette sempre p radici, dove si conta ogni radice un numero di volte pari alla sua molteplicità (ad esempio, l'equazione $x^2 - 2x + 1 = 0 = (x - 1)^2$ ha due radici, perchè si conta due volte la radice doppia $x = 1$). Queste radici sono però in generale numeri complessi (in questo caso, dato che i coefficienti del polinomio sono reali, se c'è una radice complessa deve esserci necessariamente anche la radice complessa coniugata).

risolutiva, comunque complicata⁷ Bisogna quindi risolvere in via approssimata l'equazione algebrica e il sistema lineare per potere utilizzare la (4.21).

Esercizio 4.8

(Per chi conosce già cosa sono i determinanti, le regole di operazione su di essi, nonché i sistemi di equazioni lineari di n equazioni in n incognite e come si risolvono).

Il sistema lineare (4.22) ha una ed una sola soluzione perché il determinante dei coefficienti non è nullo. Vogliamo dimostrare che questa affermazione è vera sotto la nostra ipotesi $\lambda_1 \neq \lambda_2 \neq \dots \neq \lambda_p$ calcolando questo determinante, noto come *determinante di Vandermonde*⁸.

Soluzione. Vogliamo dimostrare la formula

$$V_p := \det \begin{pmatrix} 1 & 1 & \dots & 1 \\ \lambda_1 & \lambda_2 & \dots & \lambda_p \\ \dots & \dots & \dots & \dots \\ \lambda_1^{p-1} & \lambda_2^{p-1} & \dots & \lambda_p^{p-1} \end{pmatrix} = \prod_{1 \leq l < k \leq p} (\lambda_k - \lambda_l) \quad (4.23)$$

(casi particolari: $V_1 = 1$, $V_2 = \lambda_2 - \lambda_1$). Questa formula implica che V_p è nullo se e solo se almeno due radici coincidono.

Osservazione 4.8

Il simbolo $\prod_{k=1}^n a_k$ è la notazione abbreviata del prodotto $a_1 \cdot a_2 \cdot \dots \cdot a_n$ degli n numeri $a_1, \dots, a_n$, e si chiama *produttoria*.

Procediamo per induzione: supponiamo vera la formula per V_{p-1} , e dimostriamola per V_p .

A tale scopo, sottraiamo dall'ultima riga la penultima moltiplicata per λ_1 ; poi sottraiamo dalla penultima la terz'ultima moltiplicata per λ_1 , e così via fino a sottrarre dalla seconda la prima moltiplicata per λ_1 . Tutte queste operazioni, ciascuna delle quali come sappiamo

⁷Ciò segue dal fatto che in generale un'equazione algebrica di grado $p > 4$ non è risolubile per radicali (teorema di Ruffini-Abel). Paolo Ruffini (Valentano 1765-Modena 1822), fu medico e matematico. Niels Henrik Abel (Stavanger 1802-Froland 1829), fu un matematico norvegese perfezionatosi a Parigi.

⁸Joseph Vandermonde (Parigi 1735-1796), chimico e matematico francese. Ebbe un ruolo di rilievo negli avvenimenti rivoluzionari.

non altera il valore del determinante, hanno l'effetto di riscrivere V_p sotto la forma

$$V_p = \det \begin{pmatrix} 1 & 1 & \dots & 1 \\ 0 & \lambda_2 - \lambda_1 & \dots & \lambda_p - \lambda_1 \\ 0 & \lambda_2^2 - \lambda_1 \lambda_2 & \dots & \lambda_p^2 - \lambda_1 \lambda_p \\ \dots & \dots & \dots & \dots \\ 0 & \lambda_2^{p-1} - \lambda_1 \lambda_2^{p-2} & \dots & \lambda_p^{p-1} - \lambda_1 \lambda_p^{p-2} \end{pmatrix}$$

da cui, sviluppando secondo gli elementi della prima colonna:

$$V_p = \det \begin{pmatrix} \lambda_2 - \lambda_1 & \dots & \lambda_p - \lambda_1 \\ \lambda_2^2 - \lambda_1 \lambda_2 & \dots & \lambda_p^2 - \lambda_1 \lambda_p \\ \dots & \dots & \dots \\ \lambda_2^{p-1} - \lambda_1 \lambda_2^{p-2} & \dots & \lambda_p^{p-1} - \lambda_1 \lambda_p^{p-2} \end{pmatrix}$$

Ora la prima colonna ammette il fattor comune $\lambda_2 - \lambda_1$, la seconda colonna il fattor comune $\lambda_3 - \lambda_1$; in generale la $j - 1$ -esima colonna il fattor comune $\lambda_j - \lambda_1$, $1 \leq j \leq p$. Pertanto, eseguendo i raccoglimenti (si ricordi che moltiplicare una riga o una colonna di un determinante per un numero equivale a moltiplicare tutto il determinante per quel numero) otteniamo

$$V_p = (\lambda_2 - \lambda_1)(\lambda_3 - \lambda_1) \cdots (\lambda_p - \lambda_1) \det \begin{pmatrix} 1 & 1 & \dots & 1 \\ \lambda_2 & \lambda_3 & \dots & \lambda_p \\ \dots & \dots & \dots & \dots \\ \lambda_2^{p-2} & \lambda_3^{p-2} & \dots & \lambda_p^{p-2} \end{pmatrix}$$

Ora il determinante che compare in quest'ultima formula è per definizione il determinante di Vandermonde di ordine $p - 1$. Pertanto

$$V_p = (\lambda_2 - \lambda_1)(\lambda_3 - \lambda_1) \cdots (\lambda_p - \lambda_1) V_{p-1}$$

Per l'ipotesi induttiva

$$V_{p-1} = \prod_{2 \leq l < k \leq p} (\lambda_k - \lambda_l)$$

Possiamo quindi concludere

$$\begin{aligned} V_p &= (\lambda_2 - \lambda_1)(\lambda_3 - \lambda_1) \cdots (\lambda_p - \lambda_1) V_{p-1} \\ &= (\lambda_2 - \lambda_1)(\lambda_3 - \lambda_1) \cdots (\lambda_p - \lambda_1) \prod_{2 \leq l < k \leq p} (\lambda_k - \lambda_l) \\ &= \prod_{1 \leq l < k \leq p} (\lambda_k - \lambda_l) \end{aligned}$$

il che dimostra la formula (4.23).

Capitolo 2

Successioni, serie, frazioni continue

Il materiale fin qui accumulato ci fa capire che è utile andare ad esaminare le proprietà delle successioni da un punto di vista più generale.

1 Generalità

Definizione 1.3 Sia a_n una successione. Allora essa si dice:

- (a) limitata se $\exists M > 0$ tale che $|a_n| \leq M, \forall n \in \mathbb{N} \cup \{0\}$;
- (b) monotona crescente (o monotona non decrescente) se $a_n \leq a_{n+1}, \forall n \in \mathbb{N} \cup \{0\}$;
(monotona) strettamente crescente se $a_n < a_{n+1}, \forall n \in \mathbb{N} \cup \{0\}$;
- (c) monotona decrescente (o monotona non crescente) se $a_n \geq a_{n+1}, \forall n \in \mathbb{N} \cup \{0\}$;
(monotona) strettamente decrescente se $a_n > a_{n+1}, \forall n \in \mathbb{N} \cup \{0\}$;
- (d) periodica di periodo N se esiste $N \in \mathbb{N}$ tale $a_{n+N} = a_n, \forall n \in \mathbb{N} \cup \{0\}$.

Esempio 1.3

- (a) La successione $a, a, a, \dots$ in cui tutti i termini sono uguali si dice *successione costante*. Essa è ovviamente limitata, e può essere considerata una successione periodica di periodo 1. Viceversa, ogni successione periodica di periodo 1 è la successione costante: infatti se $a_{n+1} = a_n, \forall n \in \mathbb{N} \cup \{0\}$ si ha $a_n = a_0, \forall n \in \mathbb{N} \cup \{0\}$.
- (b) La successione $a_n = \lambda^n$ è limitata se $|\lambda| \leq 1$ perché in tal caso $|\lambda^n| = |\lambda|^n \leq 1$. La successione a_n è strettamente crescente se $\lambda > 1$, strettamente decrescente se $0 < \lambda < 1$, costante se $\lambda = 1$.

- (c) La successione $a_n = (n+1)^\alpha$ è limitata se $\alpha \leq 0$, perché $|(n+1)^\alpha| \leq 1 \ \forall n \in \mathbb{N} \cup \{0\}$; non è limitata se $\alpha > 0$. a_n è inoltre strettamente crescente se $\alpha > 0$ e strettamente decrescente se $\alpha < 0$. Infatti, sia $\alpha > 0$. Allora $a_{n+1} = (n+2)^\alpha > (n+1)^\alpha = a_n$ perché $n+2 > n+1$. Allo stesso modo, se $\alpha < 0$ la successione è strettamente decrescente perché $1/(n+2) < 1/(n+1)$.
- (d) Ognuna delle successioni precedenti è monotona.
- (e) La successione $0, a, 0, a, \dots$ è periodica di periodo 2, la successione $a, b, c, a, b, c, \dots$ è periodica di periodo 3, ecc. La successione $a_n = \sin(2\pi n/N)$ è periodica di periodo N . Infatti $a_{n+N} = \sin(2\pi(n+N)/N) = \sin(2\pi n/N + 2\pi) = \sin(2\pi n/N) = a_n$.

Proviamo ora l'affermazione seguente:

Ogni successione periodica è limitata.

Infatti i soli elementi distinti di una successione periodica di periodo N sono gli N numeri $a_0, \dots, a_{N-1}$. Quindi ovviamente esiste la costante M richiesta dalla definizione. Basta sceglierla maggiore degli N numeri $|a_0|, \dots, |a_{N-1}|$.

- (f) La legge dell'incremento esponenziale genera successioni strettamente crescenti; quella del decremento esponenziale, successioni strettamente decrescenti. La relazione di ricorrenza di Fibonacci genera una successione strettamente crescente; la relazione $a_{n+2} + a_n = 0$ genera successioni limitate ma non monotone come abbiamo visto nell'osservazione 4.6 del Capitolo precedente.
- (g) Altri esempi di successioni monotone, o crescenti o decrescenti, limitate, sono i seguenti: la successione $a_{n+1} = 1 - \frac{1}{(n+1)^2}$ è strettamente crescente, perché $a_{n+1} - a_n = 1 - \frac{1}{(n+2)^2} - 1 + \frac{1}{(n+1)^2} = \frac{1}{(n+1)^2} - \frac{1}{(n+2)^2} > 0$ ed è limitata perché $0 \leq a_n \leq 1 \ \forall n$. Allo stesso modo, la successione $a_n = \frac{1}{(n+1)^2}$ è strettamente decrescente, perché $a_{n+1} - a_n = \frac{1}{(n+2)^2} - \frac{1}{(n+1)^2} = \frac{1}{(n+2)^2} - \frac{1}{(n+1)^2} < 0$ ed è limitata perché $0 \leq a_n \leq 1 \ \forall n$.

2 Limite di una successione

La nozione fondamentale per le successioni è quella di limite¹. Prima di ricordarne la definizione, facciamo un esempio.

Esempio 2.4

Consideriamo la successione $a_n = \frac{n+1}{n+2}$. Si ha:

$$a_0 = \frac{1}{2}, a_1 = \frac{2}{3}, a_2 = \frac{3}{4}, \dots, a_8 = \frac{9}{10} = 0.9, \dots, a_{98} = \frac{99}{100} = 0.99, \dots, a_{998} = \frac{999}{1000} = 0.999, \dots$$

Dunque al crescere di n si osserva che i valori della successione si avvicinano a 1, pur se il valore 1 non viene assunto per alcun n dato che l'uguaglianza $n+1 = n+2$ non ammette soluzione.

Consideriamo l'ulteriore esempio della successione $a_n = \frac{1}{n}$, $n \in \mathbb{N}$. Si ha:

$$a_1 = 1, a_2 = \frac{1}{2}, \dots, a_{10} = \frac{1}{10} = 0.1, \dots, a_{100} = \frac{1}{100} = 0.01, \dots, a_{1000} = \frac{1}{1000} = 0.001, \dots$$

Al crescere di n i valori di questa successione si avvicinano a 0, pur se, ancora, il valore 0 non viene assunto per alcun n dato che l'uguaglianza $\frac{1}{n} = 0$ non può sussistere. La nozione di limite formalizza questo tipo di fenomeno.

Definizione 2.4 Data la successione a_n e il numero $l \in \mathbb{R}$, si dice che a_n tende al limite l per $n \rightarrow \infty$, e si scrive

$$\lim_{n \rightarrow \infty} a_n = l \quad (2.1)$$

se, $\forall \epsilon > 0$, $\exists N(\epsilon) \in \mathbb{N}$ tale che $\forall n > N(\epsilon)$ si ha $|a_n - l| < \epsilon$ o, equivalentemente, $l - \epsilon < a_n < l + \epsilon$.

Osservazione 2.9

1. Equivalentemente, si dice che la successione a_n converge a l , o che ha *limite finito* l .

Una successione che ammette limite finito $l \in \mathbb{R}$ si dice *convergente*.

¹Per chi conosce già la nozione di limite per le funzioni numeriche: la nozione è la medesima. Per le successioni si può definire solo il limite per $n \rightarrow \infty$ perché l'infinito è il solo "punto di accumulazione" dei numeri naturali. La nozione di limite è stata isolata da Augustin Louis Cauchy (Parigi 1789-Parigi 1857), che fu Professore all'École Polytechnique, poi all'Università di Torino nel 1831, rientrando infine a Parigi nel 1840.

2. La condizione $|a_n - l| < \epsilon$ equivale a $l - \epsilon < a_n < l + \epsilon$ per $n > N(\epsilon)$. Se riportiamo nel piano cartesiano (x, y) i punti di ascissa n e di ordinata a_n , questa condizione significa che per *tutti* gli $n > N(\epsilon)$ i valori a_n delle ordinate si trovano nella striscia di ampiezza 2ϵ attorno alla retta $y = l$ parallela all'asse delle x . Al diminuire di ϵ l'ampiezza diminuisce, e in generale aumenta N : tuttavia la striscia conterrà sempre gli infiniti punti da $N(\epsilon)$ in poi.
3. Per verificare se una successione a_n soddisfa le condizioni della definizione 2.4 e quindi ammette limite $l \in \mathbb{R}$ in ogni caso particolare assegnato si procede nel modo seguente. Anzitutto si deve fare vedere che fra tutti gli $n \in \mathbb{N}$ che risolvono la disequazione $|a_n - l| < \epsilon$ (o, equivalentemente, del sistema $\begin{cases} a_n < l + \epsilon \\ a_n > l - \epsilon \end{cases}$) ci sono *tutti* i naturali maggiori di un opportuno $N(\epsilon)$. Quindi si cerca di risolvere il sistema e si determina $N(\epsilon)$.
4. Dimostriamo la seguente affermazione (teorema dell'unicità del limite):

Se una successione a_n ammette limite finito l , esso è unico.

In altre parole: *una successione convergente non può ammettere due limiti distinti.* Infatti se la successione a_n ammettesse i due limiti l e m con $m \neq l$, per definizione dato $\epsilon > 0$ deve esistere $N(\epsilon)$ tale che per $n > N(\epsilon)$ sono verificate entrambe le disuguaglianze $|a_n - l| < \epsilon$ e $|a_n - m| < \epsilon$. Pertanto:

$$|m - l| = |m - a_n + a_n - l| \leq |a_n - l| + |a_n - m| < 2\epsilon$$

Ora ϵ è arbitrario, e quindi così sarà 2ϵ . Dunque $l = m$ contrariamente all'ipotesi.

Esempio 2.5

1. Dimostriamo che $\lim_{n \rightarrow \infty} a_n = 0$, $a_n := \frac{1}{n^\alpha}$, se $\alpha > 0$ usando le considerazioni del punto 3 precedente. Dobbiamo fare vedere che è soddisfatta la Definizione 2.4 cioè, in questo caso:

$$\forall \epsilon > 0 \exists N(\epsilon) : \forall n > N(\epsilon) \quad \text{si ha} \quad \left| \frac{1}{n^\alpha} - 0 \right| < \epsilon, \quad \text{cioè} \quad \frac{1}{n^\alpha} < \epsilon.$$

Il dato del problema è $\epsilon > 0$; l'incognita da determinare è $N(\epsilon)$. Determiniamo $N(\epsilon)$ risolvendo la disuguaglianza $\frac{1}{n^\alpha} < \epsilon$. Questa disuguaglianza equivale a $n^\alpha > \frac{1}{\epsilon}$ da cui

$n > \frac{1}{\epsilon^{1/\alpha}}$. Quindi basta prendere $N(\epsilon) > \frac{1}{\epsilon^{1/\alpha}}$, ad esempio $N(\epsilon) = \lfloor \frac{1}{\epsilon^{1/\alpha}} \rfloor + 1$ per avere soddisfatta la definizione $\forall n > N(\epsilon)$ ($[x]$ denota la parte intera di $x \in \mathbb{R}$). Si osservi ancora che la proprietà $|a_n| < \epsilon$ sussiste per *tutti* gli $n > N(\epsilon)$. Detto qualitativamente, la successione si mantiene definitivamente più piccola di qualsiasi numero positivo, e pertanto il suo limite è zero. In generale, dire che $\lim_{n \rightarrow \infty} a_n = l$ equivale a dire che la differenza $|a_n - l|$ si mantiene *definitivamente* inferiore a qualsiasi numero positivo.

2. La successione costante, $a_n = a \forall n$ ha limite a . Infatti $|a_n - a| = 0 < \epsilon \forall \epsilon > 0$ e $\forall n \in \mathbb{N}$.
3. Un esempio numerico può servire a chiarire meglio il concetto. Sia $\alpha = 1$. Allora dalla formula precedente vediamo che se scegliamo, ad esempio, $\epsilon = 10^{-4}$ allora $N(\epsilon)$ dovrà essere almeno 10^4 ; se scegliamo $\epsilon = 10^{-8}$, allora $N(\epsilon)$ dovrà essere almeno 10^8 , e così via. L'importante è che, comunque piccolo si prenda ϵ , si possa determinare $N(\epsilon)$ in modo tale che *tutti gli infiniti termini* a_n con $n > N(\epsilon)$ si mantengano compresi fra $l - \epsilon$ e $l + \epsilon$.
4. Sia $\alpha = 2$. Allora per $\epsilon = 10^{-4}$ avremo $N(\epsilon) = 10^2$, per $\epsilon = 10^{-8}$ $N(\epsilon) = 10^4$. Si vede quindi che il limite viene raggiunto prima; in altre parole la successione $\frac{1}{n^\alpha}$ converge tanto più "rapidamente" a zero al crescere di α . in altre parole: $N(\epsilon)$ decresce al crescere di α , cosicché i valori della successione cominciano ad avvicinarsi al limite per meno di ϵ per valori più piccoli di n .
5. La successione $a_n := \lambda^n$ converge a 0 se $|\lambda| < 1$: $\lim_{n \rightarrow \infty} \lambda^n = 0$ se $|\lambda| < 1$. Per la verifica, dato $\epsilon > 0$ determiniamo $N(\epsilon)$ in modo che $\forall n > N(\epsilon)$ si abbia $|\lambda|^n < \epsilon$. Ricordiamo le seguenti proprietà dei logaritmi:

$$\begin{cases} \log_a x > \log_a y \iff x > y & \text{se } a > 1 \\ \log_a x < \log_a y \iff x > y & \text{se } a < 1 \end{cases}$$

Inoltre:

$$\log_a a^x = x, \forall x \in \mathbb{R}; \quad a^{\log_a x} = x \quad \forall x > 0$$

(cioè $\log_a x$ è la funzione inversa di a^x). Dunque la condizione $|\lambda|^n < \epsilon$ equivale a $\log_{|\lambda|} |\lambda|^n > \log_{|\lambda|} \epsilon$ da cui $n > \log_{|\lambda|} \epsilon$. Basta quindi prendere $N(\epsilon) > \log_{|\lambda|} \epsilon$. Per di più, poiché $\log_{|\lambda|} \epsilon < \frac{1}{\epsilon^{1/\alpha}}$ per ogni $\alpha < 0$, ragionando come sopra possiamo concludere che

l'esponenziale con base minore di 1 converge a zero più rapidamente di ogni potenza inversa $n^{-\alpha}$ per $n \rightarrow \infty$.

Osservazione 2.10 (*I concetti di successione limitata e di successione che ha limite finito sono profondamente diversi!*)

1. Una successione che ammette limite finito è limitata. Infatti, se $\lim_{n \rightarrow \infty} a_n = l$, per la definizione stessa di limite risulterà $|a_n - l| < \epsilon$ se $n > N(\epsilon)$. Ciò implica $||a_n| - |l|| < |a_n - l| < \epsilon$ da cui $|a_n| \leq |l| + \epsilon \forall n > N(\epsilon)$. Sia ora $\epsilon < |l|$, e poniamo $R := \max_{0 \leq n \leq N(\epsilon)} |a_n|$. Se si prende una costante M maggiore del massimo fra R e $2|l|$ si può concludere che $|a_n| < M \forall n$ e quindi la successione è limitata.
2. *Il viceversa non è vero!* Una successione può benissimo essere limitata senza ammettere limite. L'esempio più semplice è la successione $0, 1, 0, 1, \dots$ cioè $a_n = 0, n = 2p$; $a_n = 1, n = 2p + 1 \forall p \in \mathbb{N}$. È evidente infatti che questa successione "oscilla" fra 0 e 1 e non potrà mai "assestarsi" ad un unico limite. Per rendere rigorosa questa affermazione, basta osservare che la differenza fra due termini consecutivi qualsiasi è sempre 1. Questo implica la non esistenza del limite; se infatti esistesse, per definizione dato ϵ dovrebbe esistere $N(\epsilon)$ tale che $|a_n - l| < \epsilon$ per *tutti* gli $n > N(\epsilon)$. Allora:

$$1 = |a_{n+1} - a_n| = |a_{n+1} - l + l - a_n| \leq |a_{n+1} - l| + |l - a_n| \leq \epsilon + \epsilon = 2\epsilon$$

relazione evidentemente impossibile perché $\epsilon > 0$ deve essere arbitrario.

Definizione 2.5

1. Si dice che la successione a_n diverge a $+\infty$, o che ha limite $+\infty$, e si scrive

$$\lim_{n \rightarrow \infty} a_n = +\infty$$

se $\forall M > 0 \exists N(M) \in \mathbb{N}$ tale che $a_n > M$ se $n > N(M)$.

2. Si dice che la successione a_n diverge a $-\infty$, o che ha limite $-\infty$, e si scrive

$$\lim_{n \rightarrow \infty} a_n = -\infty$$

se $\forall M > 0 \exists N(M) \in \mathbb{N}$ tale che $a_n < -M$ se $n > N(M)$.

3. Una successione che soddisfa 1 o 2 si dice divergente.
4. Se una successione a_n non è nè convergente nè divergente si dice che non ammette limite.

Osservazione 2.11

1. Così come per la definizione 2.4, è utile rappresentare i punti della successione a_n nel piano cartesiano riportando in ascissa i naturali n e in ordinata i corrispondenti valori a_n . Allora l'affermazione $\lim_{n \rightarrow \infty} a_n = +\infty$ significa che per $n > N(M)$ tutti i valori a_n si mantengono superiori alla quota M . Analogamente, staranno sotto la quota $-M$ se $\lim_{n \rightarrow \infty} a_n = -\infty$.
2. Il teorema sull'unicità del limite vale ancora se si includono anche $+\infty$ e $-\infty$ fra i valori che i limiti possono assumere.
3. Anche qui la verifica consiste nel risolvere la disequazione $a_n > M$ (o $a_n < -M$) e far vedere che le soluzioni contengono tutti gli n superiori a un certo valore, che sarà il valore $N(M)$ cercato.

Esempio 2.6

1. La successione $a_n = n^\alpha$ tende a $+\infty$ se $\alpha > 0$. Basta procedere come nell'esempio precedente: scelto M a piacere, la condizione $n^\alpha > M$ implica $n > M^{1/\alpha}$ e quindi basterà prendere $N(M) \geq [M^{1/\alpha}] + 1$. Ripetendo poi il ragionamento fatto per la successione $(n+1)^{-\alpha}$, si conclude subito che n^α diverge più rapidamente al crescere di α . La divergenza della successione $a_n = n^\alpha$ si dice divergenza di potenza di esponente α .
2. La successione dell'incremento esponenziale $a_n = \lambda^n$ è divergente se $\lambda > 1$. Qui infatti $\lambda^n > M$ equivale a $n > \log_\lambda M$ e quindi basta scegliere $N(M) \geq [\log_\lambda M] + 1$ per avere $|a_n| > M \forall n > N(M)$. Anche qui, la divergenza esponenziale è più veloce di quella di qualsiasi potenza.

3. La successione $a_n = \lambda^n$ non ammette limite se $\lambda < -1$. Infatti si ha

$$a_{2n} = \lambda^{2n} \rightarrow \infty; \quad a_{2n+1} = \lambda \lambda^{2n} \rightarrow -\infty$$

perché $\lambda^{2n} > 0$ mentre $\lambda < 0$, e ciò basta a provare la non esistenza del limite.

3 Alcuni teoremi fondamentali sui limiti

Enunciamo, proponendo la dimostrazione come esercizio, alcune delle proprietà principali dell'operazione di limite.

Teorema 3.1 *Siano a_n e b_n due successioni convergenti, con $\lim_{n \rightarrow \infty} a_n = a$, $\lim_{n \rightarrow \infty} b_n = b$.*

Allora

1. *Il limite della somma vale la somma dei limiti:*

$$\lim_{n \rightarrow \infty} (a_n + b_n) = a + b$$

2. *Il limite del prodotto vale il prodotto dei limiti:*

$$\lim_{n \rightarrow \infty} (a_n \cdot b_n) = a \cdot b$$

In particolare se b_n è una costante, $b_n = b \forall n$, $\lim_{n \rightarrow \infty} (ba_n) = b \lim_{n \rightarrow \infty} (a_n) = ab$.

3. *Se $b \neq 0$ il limite del quoziente vale il quoziente dei limiti:*

$$\lim_{n \rightarrow \infty} \frac{a_n}{b_n} = \frac{a}{b}$$

Sia ora a_n una successione limitata, $|a_n| \leq M \forall n \in \mathbb{N}$, e sia b_n divergente a $+\infty$ o a $-\infty$.

Allora:

1. *La successione $a_n + b_n$ è divergente a $+\infty$ o a $-\infty$ rispettivamente.*

2. *Se $\lim_{n \rightarrow \infty} a_n = a \neq 0$ la successione $a_n \cdot b_n$ è divergente. Più precisamente: $+\infty$ ($-\infty$) se $a > 0$ ($a < 0$). (4 casi a seconda della regola dei segni).*

3. *La successione $\frac{a_n}{b_n}$ converge a zero: $\lim_{n \rightarrow \infty} \frac{a_n}{b_n} = 0$.*

Esercizio 3.9

Dimostrare il teorema.

Soluzione. Seguire il corso di Analisi.

Esempio 3.7

1. La successione $a_n = 1 + \frac{1}{n^2} + 2^{-n}$ converge a 1. Infatti $\lim_{n \rightarrow \infty} \frac{1}{n^2} = 0$ e $\lim_{n \rightarrow \infty} 2^{-n} = 0$.
2. La successione $c_n = [1 + \frac{1}{n^2}] \cdot 2^{-n}$ converge a 0. Infatti anzitutto la successione $a_n = [1 + \frac{1}{n^2}]$ converge a 1. Questo perché $\frac{1}{n^2}$ converge a 0, la successione 1 è costante, e il limite della somma è la somma dei limiti; poi la successione 2^{-n} converge a 0, e quindi possiamo applicare la proprietà che il limite del prodotto è il prodotto dei limiti.
3. La successione $a_n = \sin(2\pi n/p) + n^2$ è divergente positivamente. Infatti la successione $\sin(2\pi n/p)$ è limitata e la successione n^2 diverge positivamente.
4. La successione $a_n = \frac{\sin(2\pi n/p)}{n^2}$ tende al limite 0 perché il numeratore è limitato e il denominatore diverge positivamente.
5. La successione $a_n = \frac{\sin(2\pi n/p)}{\lambda^n}$ è convergente a 0 se $\lambda > 1$ perché la successione $a_n = \sin(2\pi n/p)$ è limitata e la successione λ^n è divergente positivamente se $\lambda > 1$.
6. Se a_n è una successione limitata qualsiasi, la successione $\frac{a_n}{\lambda^n}$ converge a 0 se $\lambda > 1$ perché la successione λ^n è divergente positivamente se $\lambda > 1$.

Data una successione, è cosa di grande rilevanza stabilire dei criteri che permettano di determinarne immediatamente la convergenza la divergenza. Un criterio fondamentale è il seguente:

Teorema 3.2 *Sia a_n una successione monotona. Allora:*

1. *Se a_n è limitata essa ammette limite, cioè esiste a tale che $\lim_{n \rightarrow \infty} a_n = a$.*
2. *Se a_n non è limitata essa diverge a $+\infty$ o a $-\infty$.*

Dimostrazione. Seguire il corso di Analisi.

Esempio 3.8 (*Il numero di Nepero*)

Consideriamo la successione il cui termine generale è

$$a_n = \left(1 + \frac{1}{n}\right)^n$$

Dimostriamo che a_n è crescente. Per la forma del binomio di Newton (2.1) con $x = \frac{1}{n}$ e

$$a_k(n) = \frac{n!}{k!(n-k)!} \text{ si ha}$$

$$\begin{aligned} a_n &= \left(1 + \frac{1}{n}\right)^n = \sum_{k=0}^n \frac{n!}{k!(n-k)!} \frac{1}{n^k} = \\ &= \sum_{k=0}^n \frac{1}{k!} \frac{n(n-1)\cdots(n-k+1)}{n^k} = \sum_{k=0}^n \frac{1}{k!} \left(1 - \frac{1}{n}\right) \left(1 - \frac{2}{n}\right) \cdots \left(1 - \frac{k-1}{n}\right) \end{aligned}$$

Se adesso andiamo a mettere $n+1$ al posto di n in questa formula il numero degli addendi, tutti positivi, aumenta e il valore di ciascuno non diminuisce. Quindi $a_n < a_{n+1}$ e la successione è strettamente crescente. Inoltre dalla formula precedente segue subito, dato che tutti i prodotti entro parentesi sono minori di 1:

$$\begin{aligned} \left(1 + \frac{1}{n}\right)^n &< \sum_{k=0}^n \frac{1}{k!} = 1 + 1 + \frac{1}{2} + \frac{1}{2 \cdot 3} + \cdots + \frac{1}{1 \cdot 2 \cdot 3 \cdots n} < \\ &1 + \left(1 + \frac{1}{2} + \frac{1}{2^2} + \cdots + \frac{1}{2^{n-1}}\right) = 1 + \frac{1 - 2^{-n}}{1 - 1/2} < 1 + 2 = 3 \end{aligned}$$

Qui abbiamo fatto uso, nel provare la seconda disuguaglianza, della constatazione che $2^k = 2 \cdots 2 < 2 \cdot 3 \cdots k = k!$ e della formula

$$1 + \frac{1}{2} + \frac{1}{2^2} + \cdots + \frac{1}{2^{n-1}} = \frac{1 - 2^{-n}}{1 - 1/2}$$

(la cui verifica è immediata: basta moltiplicare ambo i membri per $1 - 1/2$). Quindi la successione è anche limitata, e pertanto convergente. Il suo limite è un numero reale fra 2 e 3, che si denota e e porta il nome di *numero di Nepero*². Concludendo:

Definizione 3.6

$$e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n$$

²Dal nome di John Napier, o Neper (Edimburgo 1550-1617), in latino Neperus, italianizzato in Nepero.

4 Serie numeriche

Una classe di successioni di estrema importanza è quella delle serie numeriche. Le serie numeriche costituiscono il procedimento per definire e calcolare somme di infiniti numeri. Ce ne occuperemo ora brevemente. Cominciamo con due esempi fondamentali.

Esempio 4.9 (*Progressioni aritmetiche e geometriche*)

1. Supponiamo di volere calcolare la somma S_n dei primi n numeri naturali: $S_n := 1 + 2 + 3 + \dots$ che possiamo indicare anche con $S_n = \sum_{k=1}^n k$. S_n è ovviamente una successione crescente, detta *progressione aritmetica di ragione 1* perché ogni termine nella somma si ottiene dal precedente aggiungendo 1. Si ha:

$$S_n = \frac{n(n+1)}{2} \quad (4.2)$$

Dimostriamo questa formula³ per induzione. Per $n = 1$ si ha $S_1 = 1$ come deve. Assumiamo allora che sia $S_{n-1} = \frac{n(n-1)}{2}$. Allora:

$$S_n = S_{n-1} + n = \frac{n(n-1)}{2} + n = \frac{n(n+1)}{2}.$$

2. Supponiamo ora di volere calcolare la somma G_n delle prime $n + 1$ potenze di un numero fissato $x \in \mathbb{R}$:

$$G_n := 1 + x + x^2 + \dots + x^n = \sum_{k=0}^n x^k$$

G_n è detta *progressione geometrica di ragione x* perché ogni termine nella somma si ottiene dal precedente moltiplicandolo per x . Anche la formula per G_n è ben nota:

$$G_n = \frac{1 - x^{n+1}}{1 - x}, x \neq 1; \quad G_n = n + 1, \quad x = 1. \quad (4.3)$$

La si deduce immediatamente dividendo per $1 - x$ il prodotto notevole $1 - x^{n+1} = (1 - x)(1 + x + x^2 + \dots + x^n)$.

³La leggenda vuole che sia stata scritta per la prima volta da Gauss nel 1787 all'età di 10 anni quando il maestro per punizione gli ingiunse di calcolare la somma di tutti i numeri da 1 a 100. Carl Friedrich Gauss (Braunschweig 1777-Göttingen 1855), fu Professore di Astronomia e poi di Matematica all'Università di Göttingen.

Il procedimento per ottenere S_n e G_n sommando consecutivamente i termini delle successioni $a_n = n$ e $a_n = x^n$ a partire dal primo ($n = 0$) si può applicare ad ogni successione a_n . La serie numerica corrispondente esamina l'andamento di questa successione per $n \rightarrow \infty$, in altre parole il limite della successione delle somme parziali.

Definizione 4.7

1. Data la successione $a_k, k = 0, 1, \dots$, dicesi somma parziale n -esima da essa dedotta, denotata s_n , la somma dei primi $n + 1$ numeri a_k :

$$s_n := a_0 + a_1 + \dots + a_n = \sum_{k=0}^n a_k \quad (4.4)$$

ossia

$$s_0 = a_0; \quad s_1 = a_0 + a_1; \quad s_2 = a_0 + a_1 + a_2, \quad s_n = a_0 + a_1 + \dots + a_n = s_{n-1} + a_n$$

La successione s_n delle somme parziali dicesi serie numerica di ridotte s_n , e si indica col simbolo $\sum_{n=0}^{\infty} a_n$ o anche $a_0 + a_1 + a_2 + \dots$.

2. Se la successione delle somme parziali ammette limite $s \in \mathbb{R}$, cioè se $\lim_{n \rightarrow \infty} s_n = s$ si dice che la serie numerica di ridotte s_n è convergente, e si scrive

$$\sum_{n=0}^{\infty} a_n = s \quad (4.5)$$

Il limite s si dice somma della serie.

3. Se la successione s_n diverge positivamente (negativamente) scriveremo $\sum_{n=0}^{\infty} a_n = +\infty$ ($\sum_{n=0}^{\infty} a_n = -\infty$) e diremo che la serie è divergente.

Esempio 4.10

La successione G_n dell'esempio precedente si dice *serie geometrica*. Dimostriamo quanto segue:

1. Se $|x| < 1$ la serie geometrica è convergente e la sua somma vale $s = \frac{1}{1-x}$. In formule:

$$\sum_{n=0}^{\infty} x^n = \frac{1}{1-x}, \quad |x| < 1$$

2. Se $x = -1$ la serie geometrica è oscillante. Se $x < -1$ la somma non esiste, nè finita nè infinita.

Per dimostrare queste affermazioni consideriamo la successione $G_n = \sum_{k=1}^n x^k = \frac{1-x^{n+1}}{1-x}$ delle somme parziali. Poiché $|x| < 1$ sappiamo (Esempio 2.5.5) che $\lim_{n \rightarrow \infty} x^{n+1} = 0$. Quindi, per l'affermazione 3 del Teorema 3.1 abbiamo

$$\sum_{n=0}^{\infty} x^n = \lim_{n \rightarrow \infty} \frac{1-x^{n+1}}{1-x} = \frac{1}{1-x}$$

e ciò prova la prima affermazione.

La verifica dell'asserzione valida per $x < -1$ segue dal fatto già dimostrato che la successione x^n non ha limite per $x < -1$. Infine se $x = 1$ $G_n = n$ e quindi la successione delle somme parziali diverge positivamente.

Rimane il caso $x = -1$. Qui si vede subito che $G_{2n} = 0$, $G_{2n+1} = 1$; quindi come è facile verificare la successione delle somme parziali pur essendo limitata non ammette limite e la serie non converge. Essa però converge a $\frac{1}{2}$ nel senso di Cesàro come vedremo nel punto della dimostrazione del Teorema 5.1.

La serie aritmetica diverge positivamente. Infatti $S_n = \frac{n(n+1)}{2}$ ed abbiamo già visto che questa successione tende a $+\infty$.

Osservazione 4.12

1. Se $a_n > 0$ ($a_n < 0$) la serie $\sum_{n=0}^{\infty} a_n$ si dice a termini positivi (negativi). La successione delle somme parziali è in questi casi monotona. Quindi per il Teorema 3.2 una serie a termini positivi (negativi) o converge o diverge positivamente (negativamente).
2. Nel caso precedente si vede subito che condizione necessaria (ma non sufficiente!) affinché la serie converga è che sia $\lim_{n \rightarrow \infty} a_n = 0$. Altrimenti la successione s_n delle somme parziali, che è monotona, non potrebbe essere limitata e dovrebbe necessariamente divergere. Ad esempio, si può dimostrare che la serie a termini positivi $\sum_{k=1}^{\infty} \frac{1}{n^\alpha}$ converge per $\alpha > 1$ mentre diverge per $0 < \alpha \leq 1$ anche se $\lim_{n \rightarrow \infty} \frac{1}{n^\alpha} = 0$. Si può dimostrare che questo risultato rimane vero anche senza assumere che la successione a_n abbia segno costante. La conclusione è che si può dimostrare il seguente:

Teorema 4.3 La serie numerica $\sum_{n=0}^{\infty} a_n$ può convergere solo se $\lim_{n \rightarrow \infty} a_n = 0$.

Questo teorema dà una condizione necessaria per la convergenza. Esistono molte condizioni sufficienti. La più utile è forse la seguente, basata sul confronto.

Anzitutto diciamo che una serie numerica $\sum_{n=0}^{\infty} a_n$ converge *assolutamente* se converge la serie dei suoi valori assoluti, cioè se

$$\sum_{n=0}^{\infty} |a_n| < +\infty$$

Date due serie numeriche $\sum_{n=0}^{\infty} a_n$, $\sum_{n=0}^{\infty} b_n$, $b_n > 0$, diremo poi che la prima è *maggiorata* dalla seconda (equivalentemente, la seconda *maggiora* la prima) se esiste $K > 0$ tale che

$$|a_n| \leq K b_n \quad \forall n \in \mathbb{N}$$

Teorema 4.4

1. Se $\sum_{n=0}^{\infty} |a_n| < +\infty$ allora anche $\sum_{n=0}^{\infty} a_n$ converge, cioè la convergenza assoluta implica la convergenza semplice (mentre il viceversa non è vero).
2. Se la serie $\sum_{n=0}^{\infty} |a_n|$ è maggiorata dalla serie $\sum_{n=0}^{\infty} b_n$, $b_n > 0$ e quest'ultima converge, anche la prima converge e

$$\sum_{n=0}^{\infty} |a_n| \leq K \sum_{n=0}^{\infty} b_n$$

Consequentemente se $\sum_{n=0}^{\infty} |a_n|$ diverge allora anche $\sum_{n=0}^{\infty} b_n$ diverge.

Dimostrazione.

Seguire il corso di Analisi.

Osservazione 4.13

Sia a_n un successione maggiorata da x^n con $0 < x < 1$: $|a_n| \leq x^n$. Allora la serie $\sum_{n=0}^{\infty} a_n$ converge assolutamente perché è maggiorata da una serie geometrica convergente: Ad esempio, la serie

$$\sum_{n=0}^{\infty} \frac{a^n}{n!}$$

detta *serie esponenziale* per il motivo chiarito nell'esercizio seguente, è assolutamente convergente $\forall a \in \mathbb{R}$. Infatti, sia $a > 0$ e al solito si denoti $[a]$ la parte intera di a . Allora

$$\sum_{n=0}^{\infty} \frac{a^n}{n!} = \sum_{n=0}^{[a]} \frac{a^n}{n!} + \sum_{n=[a]+1}^{\infty} \frac{a^n}{n!}$$

Il primo termine è una somma finita, e quindi occorre e basta mostrare la convergenza del secondo. Si ha anzitutto:

$$\sum_{n=[a]+1}^{\infty} \frac{a^n}{n!} = a^{[a]} \sum_{k=1}^{\infty} \frac{a^k}{([a]+k)!}$$

D'altra parte:

$$\frac{a^k}{([a]+k)!} \leq \frac{a^k}{([a]+1)^k} = b^k, \quad b := \frac{a}{[a]+1} < 1$$

Pertanto:

$$\sum_{k=1}^{\infty} \frac{a^k}{([a]+k)!} \leq \sum_{k=1}^{\infty} b^k = b \sum_{k=0}^{\infty} b^k = \frac{b}{1-b}.$$

se $a > 0$, ma questo ovviamente implica la convergenza assoluta per ogni $a \in \mathbb{R}$ perché $\frac{|-a|^n}{n!} = \frac{a^n}{n!}$ con $a > 0$.

Esercizio 4.10 (*Per chi conosce il calcolo infinitesimale*)

Dimostrare che $\sum_{n=0}^{\infty} \frac{a^n}{n!} = e^a$, $\forall a \in \mathbb{R}$.

Suggerimento. Ammettendo che $\frac{d}{da} \sum_{n=0}^{\infty} \frac{a^n}{n!} = \sum_{n=0}^{\infty} \frac{d}{da} \frac{a^n}{n!}$ far vedere che

$$\frac{d}{da} \sum_{n=0}^{\infty} \frac{a^n}{n!} = \sum_{n=0}^{\infty} \frac{a^n}{n!}.$$

Esercizio 4.11

Si dimostra in Analisi che la serie $\sum_{n=1}^{\infty} \frac{1}{n^\alpha}$ (detta *serie armonica generalizzata* per $\alpha \neq 1$; *serie armonica* per $\alpha = 1$) è convergente se $\alpha > 1$, mentre diverge se $\alpha \leq 1$. Far vedere che qualsiasi serie $\sum_{n=0}^{\infty} a_n$ tale che per $n \in \mathbb{N}$ si abbia $|a_n| \leq \frac{K}{n^\alpha}$ con $\alpha > 1$ è convergente, mentre è divergente se $a_n \geq \frac{M}{n^\alpha}$ con $\alpha \leq 1$. Qui K e M sono costanti positive.

Concludiamo questo paragrafo con un'applicazione delle serie numeriche alla filosofia.

Esempio 4.11 (*Paradossi di Zenone*)

Volendo mettere in evidenza che le affermazioni dedotte dal puro ragionamento possono risultare in contraddizione con le osservazioni empiriche, il filosofo sofista Zenone di Elea (V secolo A.C.; Elea era una città della Magna Grecia vicino a Metaponto, in Lucania) enunciò un paradosso diventato celebre che riporteremo qui nelle sue due varianti più conosciute⁴.

1. *La freccia scoccata dall'arco non colpirà mai il bersaglio, o paradosso della dicotomia.*

Prima di colpire il bersaglio la freccia dovrà percorrere la metà della distanza che separa quest'ultimo dall'arco, e poi la metà della metà restante, e così via. Percorrere questi infiniti tratti richiederà ancora un tempo infinito, e quindi la freccia non colpirà mai il bersaglio.

2. *Il piè veloce Achille non raggiungerà mai la tartaruga.*

Achille, il più veloce eroe omerico⁵, si lancia all'inseguimento di una tartaruga partita in vantaggio su di lui, alla distanza d . Ebbene non la raggiungerà mai: infatti il piè veloce impiegherà un certo tempo per raggiungere la posizione iniziale d ; nel frattempo la tartaruga avrà percorso un certo tratto e si sarà portata nella posizione d' , che Achille impiegherà un certo tempo a raggiungere, ecc.; quindi Achille non potrà mai raggiungere la tartaruga.

Come la matematica risolve il paradosso della dicotomia. Assumiamo per comodità che la distanza fra l'arco e il bersaglio valga 1. Esplicitando l'assunzione implicita nella formulazione del paradosso che la velocità di percorrenza della freccia sia costante, avremo che, fatto $\frac{t}{2}$ il tempo necessario a percorrere la prima metà della distanza, il tempo necessario a percorrere la metà della metà sarà $\frac{t}{4}$, e così via. Il tempo T necessario per percorrere l'intera distanza sarà quindi dato da

$$T = t \left(\frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \dots \right)$$

Si tratta di una serie geometrica di ragione $\frac{1}{2}$ e punto iniziale $\frac{1}{2}$. Pertanto, raccogliendo $\frac{1}{2}$ a fattor comune:

$$T = \frac{t}{2} \sum_{n=0}^{\infty} \left(\frac{1}{2} \right)^n = \frac{t}{2} \frac{1}{1 - \frac{1}{2}} = t$$

Quindi il mobile impiegherà, per giungere al punto di arrivo, esattamente il doppio del tempo che ha impiegato a percorrere la metà della distanza che separa il punto di partenza

⁴Tramandateci da Aristotele nel Libro VI della *Fisica*.

⁵Nell'Iliade si narra come Achille raggiunse Ettore dopo averlo inseguito per tre giri di corsa attorno alle mura di Troia. Ettore fuggiva l'ira funesta di Achille perché gli aveva ucciso l'amico Patroclo.

da quello di arrivo, così come ci suggerisce la nostra intuizione. Dunque il paradosso si risolve semplicemente sommando una serie geometrica di ragione $\frac{1}{2}$.

Esercizio 4.12

Risolvere il paradosso di Achille e della tartaruga, dimostrando che il raggiungimento avrà luogo se egli è K volte più veloce della tartaruga con $K > 1$.

La morale del tutto è: *Non è affatto detto che la somma di infiniti numeri debba essere infinita. Può benissimo essere finita!*

5 Successioni di 0 e 1

Le successioni in cui i numeri a_n assumono i soli valori 0 e 1, e cioè le successioni del tipo $0, 1, 1, 0, 0, 0, 1, 0, \dots$ si incontrano in molti casi significativi. Vediamone subito alcuni particolarmente importanti; altri li vedremo nei capitoli successivi.

Esempio 5.12 (La numerazione in base 2)

La numerazione in base 10, o numerazione decimale, significa che possiamo scrivere ogni intero in uno ed un solo modo sotto la forma $n = x_{N-1}x_{N-2} \dots x_0$ dove x_i prendono i valori $0, 1, \dots, 9$ (cifre). Qui x_0 indica il numero delle unità, x_1 quello delle decine, e così via fino a x_N che indica il numero delle potenze di 10 di ordine N :

$$n = x_N 10^N + x_{N-1} 10^{N-1} + \dots + x_1 10^1 + x_0$$

Ad esempio 13045 significa 1 decina di migliaia, 3 migliaia, 0 centinaia, 3 decine e 5 unità:

$$13045 = 1 \times 10^4 + 3 \times 10^3 + 0 \times 10^2 + 4 \times 10 + 5 \times 1$$

cioè $x_0 = 5, x_1 = 4, x_2 = 0, x_3 = 3, x_4 = 1$. Allo stesso modo, la numerazione in base 2 significa scrivere ogni intero sotto la forma

$$n = y_N y_{N-1} \dots y_0$$

dove però, essendoci solo due "cifre" le y_i prendono solo i valori 0 e 1. In tal caso y_0 rappresenta le potenze 2^0 cioè l'unità, y_1 la potenza di grado 1 cioè 2, ecc. In generale, come sopra,

$$n = y_N 2^N + y_{N-1} 2^{N-1} + \dots + y_1 2 + y_0.$$

Calcoliamo ad esempio la decomposizione binaria di 13405: la potenza più alta di 2 contenuta in 13405 è $2^{13} = 8192$. La potenza di 2 più alta contenuta in $13405 - 8192 = 5213$ è $2^{12} = 4096$; la potenza di 2 più alta contenuta in $5213 - 4096 = 1117$ è $2^{10} = 1024$; la potenza di 2 più alta contenuta in $1117 - 1024 = 93$ è $2^6 = 64$; la potenza di 2 più alta contenuta in $93 - 64 = 29$ è $2^4 = 16$; la potenza di 2 più alta contenuta in $29 - 16 = 13$ è $2^3 = 8$; la potenza di 2 più alta contenuta in $13 - 8 = 5$ è $2^2 = 4$; la potenza di 2 più alta contenuta in $5 - 4 = 1$ è $2^0 = 1$. Dunque sono presenti solo le potenze di 2 di esponente 13, 12, 10, 6, 4, 3, 2, 0. Quindi

$$13405 = 2^{13} + 2^{12} + 0 \times 2^{11} + 2^{10} + 0 \times 2^9 + 0 \times 2^8 + 0 \times 2^7 + 2^6 + 0 \times 2^5 + 2^4 + 2^3 + 2^2 + 0 \times 2^1 + 1$$

In altre parole, $13405 = 11010001011101$ in base 2.

Esercizio 5.13

Calcolare le decomposizioni binarie di 11457 e 18951.

Risposta. $11457 = 10110011000001$, $18951 = 100101000000111$.

Esercizio 5.14

Se la numerazione binaria è uno strumento così comodo, perché non lo si adotta nell'uso quotidiano invece che lasciarla ai calcolatori?

La decomposizione binaria di un numero naturale genera una successione *finita*, detta anche *stringa*, i cui elementi assumono solo i valori 0, 1. Come tale, essa genera la decomposizione binaria della parte intera di qualsiasi numero reale. Successioni infinite corrispondono alla decomposizione binaria della parte decimale, che si sviluppa in modo analogo come passiamo ora a ricordare. Trattandosi della parte decimale, possiamo prendere x fra 0 e 1: $0 \leq x \leq 1$. Consideriamo la sua decomposizione decimale, $x = 0.x_1x_2\dots$, $0 \leq x_i \leq 9$. Essa corrisponde a rappresentare x nella forma seguente:

$$x = x_1 10^{-1} + x_2 10^{-2} + x_3 10^{-3} + \dots \quad (5.6)$$

dove x_1 indica quanti decimi ci sono, x_2 quanti centesimi, ecc. nel senso che si può dimostrare che la serie a secondo membro converge a x . Esempio:

$$x = 0.3709 : x_1 = 3, x_2 = 7, x_3 = 0, x_4 = 9, x_k = 0, k \geq 5.$$

Possiamo quindi mettere in corrispondenza ad ogni numero reale $0 \leq x \leq 1$ la successione delle cifre del suo sviluppo decimale:

$$x \mapsto r_n := (x_1, x_2, x_3, \dots) \quad (5.7)$$

dove ciascun elemento x_k è un intero fra 0 e 9 (compresi). Riguardo questa corrispondenza richiamiamo alcuni fatti ben noti:

1. A priori ad un medesimo numero x possono corrispondere due successioni distinte: ad esempio il numero 1 può essere scritto come $1,000\dots$ o equivalentemente come $0,9999\dots$; allo stesso modo le rappresentazioni equivalenti $0.5 = 0.4999\dots$ generano le due successioni differenti $5,0,0,\dots$ e $4,9,9,\dots$. Inversamente, le successioni del tipo $x_1, x_2, \dots, x_k, 9, \dots, 9, \dots$ e le successioni del $x_1, x_2, \dots, x_k + 1, 0, 0, \dots$ generano lo stesso numero.
2. Se da un certo indice m *geq* 1 in poi la successione $(x_1, x_2, x_3, \dots)$ è periodica, cioè se la successione $(x_m, x_{m+1}, x_{m+2}, \dots)$ è periodica, il numero corrispondente x è razionale e viceversa. Ad esempio $3/8 = 0.37500000\dots = 0.3749999999\dots$

Allo stesso modo, dato $0 \leq x \leq 1$, la sua *rappresentazione binaria* $x_1, x_2, x_3, \dots$ corrisponde a scrivere x nel modo seguente:

$$x = x_1 2^{-1} + x_2 2^{-2} + x_3 2^{-3} + x_4 2^{-4} + \dots \quad (5.8)$$

dove $x_i = 0$ o $x_i = 1$. Come sopra, si usa l'abbreviazione $x = (x_1, x_2, x_3, \dots)$. Anche qui non è detto che lo sviluppo binario di un numero sia unico. Ad esempio, il numero 0.5 ammette in base 2 le due rappresentazioni $1,0,0,\dots$ e $0,1,1,1,\dots$. Infatti inserendo la seconda rappresentazione nella (5.8) troviamo

$$x = 0 \times 2^{-1} + \lim_{n \rightarrow \infty} \sum_{k=2}^n 2^{-k} = \frac{1}{4} \lim_{n \rightarrow \infty} \sum_{k=0}^{n-2} 2^{-k} = \frac{1}{4} \frac{1}{1 - 1/2} = \frac{1}{2}$$

(Qui abbiamo usato, al solito, la formula $\lim_{n \rightarrow \infty} \sum_{k=0}^n 2^{-k} = \frac{1}{1 - \frac{1}{2}} = 2$ per la somma della serie geometrica di ragione $\frac{1}{2}$). Come nel caso decimale, a numeri razionali corrispondono rappresentazioni binarie periodiche.

Esercizio 5.15

1. Il numero razionale $\frac{1}{3}$ ha la rappresentazione decimale $(3, 3, 3, \dots) = 0.333\dots$. Trovare la sua rappresentazione binaria.

Soluzione Poiché $\frac{1}{3} < \frac{1}{2}$, avremo $x_1 = 0$. x_2 sarà quindi definito dalla condizione $\frac{1}{3} = \frac{1}{4} + y_2$ da cui $x_2 = 1$ e $y_2 = \frac{1}{3} - \frac{1}{4} = \frac{1}{12}$. Proseguendo, $x_3 = 0$ perché $\frac{1}{12} < \frac{1}{8}$, $\frac{1}{12} = \frac{1}{16} + y_4$ da cui $x_4 = 1$, $y_4 = \frac{1}{12} - \frac{1}{16} = \frac{1}{48}$. Procedendo ancora, si trova subito $x_{2n+1} = 0$, $x_{2n} = 1$, $y_{2n} = \frac{1}{3 \cdot 2^{2n}}$. In conclusione, l'espansione binaria di $\frac{1}{3}$ è la successione $(0, 1, 0, 1, \dots)$, ovvero

$$\frac{1}{3} = \frac{1}{4} + \frac{1}{16} + \frac{1}{64} + \frac{1}{2^{2n}} + \dots = \frac{1}{4} \left(1 + \frac{1}{4} + \frac{1}{16} + \dots \right)$$

Infatti la somma entro parentesi è una serie geometrica di ragione $\frac{1}{4}$. Essa pertanto vale $\frac{1}{1 - 1/4}$. Quindi

$$\frac{1}{4} + \frac{1}{16} + \frac{1}{64} + \frac{1}{2^{2n}} + \dots = \frac{1}{4} \frac{1}{1 - 1/4} = \frac{1}{3}$$

2. Trovare gli sviluppi binari delle frazioni: $\frac{1}{5}, \frac{1}{7}, \frac{1}{11}$.

Soluzione. Consideriamo $\frac{1}{5}$. Si ha $\frac{1}{5} < \frac{1}{2}$, $\frac{1}{5} < \frac{1}{4}$. Quindi i primi due termini dello sviluppo binario sono 00. Per trovare i successivi, procediamo come sopra: poniamo $\frac{1}{5} = \frac{1}{8} + y_3$ da cui $y_3 = \frac{3}{40}$; $\frac{3}{40} = \frac{1}{16} + y_4$ da cui $y_4 = \frac{1}{80}$; poi $\frac{1}{80} < \frac{1}{32}$ e $\frac{1}{80} < \frac{1}{64}$ da cui $y_5 = y_6 = 0$; proseguendo si ha $\frac{1}{80} = \frac{1}{128} + y_7$ da cui $y_7 = \frac{1}{16}(\frac{1}{5} - \frac{1}{8}) = \frac{1}{16}y_3$. Proseguendo ancora, si trova $y_{n+k} = \frac{1}{16}y_3$ da cui la periodicità di periodo 4. Concludendo, lo sviluppo di $\frac{1}{5}$ in base 2 è:

$$\frac{1}{5} = (001100110011\dots).$$

che possiamo scrivere come $\frac{1}{5}(x_1, x_2, x_2, \dots, x_n \dots)$ con $x_{4k-1} = x_{4k} = 1$; $x_{4k-2} = x_{4k-3} = 0, k = 1, 2, \dots$. Per verificare il risultato, sommiamo ancora le serie geometriche. Si ha:

$$\sum_{n=1}^{\infty} \frac{x^n}{2^n} = \sum_{k=1}^{\infty} \frac{1}{2^{4k-1}} + \sum_{k=1}^{\infty} \frac{1}{2^{4k}} = \sum_{k=1}^{\infty} \frac{2}{2^{4k}} + \sum_{k=1}^{\infty} \frac{1}{2^{4k}} =$$

$$\sum_{k=1}^{\infty} \frac{3}{2^{4k}} = \frac{3}{16} \cdot \sum_{k=0}^{\infty} \frac{1}{2^{4k}} = \frac{3}{16} \frac{1}{1 - \frac{1}{16}} = \frac{1}{5}$$

Esercizio 5.16

1. Trovare lo sviluppo *ternario* di $\frac{1}{3}, \frac{2}{9}, \frac{1}{27}, \frac{18}{81}$.

Risposta: 10000..., 020000..., 00100000..., 0200000....

2. Trovare lo sviluppo ternario di $\frac{1}{4}$.

Risposta: 020202...

Osservazione 5.14

La coppia $(0, 1)$ definisce l'unità elementare di informazione (vero-falso, acceso-spento, ecc.) e viene chiamata *bit*⁶. Mettendo insieme n bits successivi si ottengono n -ple (o *stringhe* di lunghezza n) ciascun elemento delle quali vale 0 o 1. Con successioni formate da n elementi ciascuno dei quali può assumere solo il valore 0 o 1, si possono formare 2^n stringhe diverse: infatti per $n = 1$ ce ne sono due, per $n = 2$ ce ne sono 4 perché ognuna delle due precedenti ne genera 2 aggiungendo uno 0 o un 1, e così via fino a n . Si possono così codificare 2^n informazioni indipendenti perché due stringhe che differiscono anche in un solo sito sono stringhe differenti (ad esempio, si possono codificare 2^n numeri diversi). Il numero n è per definizione il numero di bits. Si definisce poi *byte* $8 = 2^3$ bit, cioè il numero di bit necessari per codificare un carattere alfanumerico ASCII (ci sono $2^8 = 256$ caratteri ASCII). Dunque con n bytes si codificano 2^n informazioni diverse. I kilobytes sono mille bytes, e i megabytes sono un milione di bytes. Dire che il processore di un calcolatore ha la memoria RAM⁷ di (ad esempio) 64 megabytes significa che esso è in grado di codificare e gestire contemporaneamente $2^{64 \cdot 1000 \cdot 1000}$ informazioni. Ogni informazione è rappresentata da una stringa lunga 64.000.000 di termini ciascuno dei quali vale 0 oppure 1.⁸ Il modo concreto di realizzare

⁶Iniziali delle tre parole inglesi *binary information unit* = unità di informazione binaria.

⁷Acronimo di *Random Access Memory*.

⁸L'invenzione dell'aritmetica binaria e l'idea di codificare tutte le informazioni per suo tramite risale a Leibnitz. Gli piacque talmente che egli credette di vedervi l'immagine della creazione. Immaginò infatti che 1 potesse rappresentare Dio, e 0 il nulla, e che l'Essere Supremo avesse creato dal nulla tutti gli esseri di questo mondo, codificati da una successione di 0 e 1 nell'aritmetica binaria.

l'alternativa 0, 1 è definito dall'hardware del calcolatore; all'inizio dell'era dei calcolatori, subito dopo la seconda guerra mondiale, si impiegavano reticoli composti da nuclei di ferrite, che si magnetizzavano immettendovi della corrente elettrica e si smagnetizzavano immediatamente interrompendola. In questo modo si realizzava un'alternativa binaria; le codificazioni venivano ottenute tramite il passaggio di corrente alternata ad alta frequenza nei vari siti del reticolo.⁹

Esercizio 5.17

Sia dato un alfabeto composto da N lettere. Definiamo *parola di lunghezza n* una successione arbitraria di n lettere estratte dall'alfabeto. Quante parole di n lettere si potranno formare?

Soluzione. Ogni parola è una successione finita di n elementi, ciascuno dei quali può assumere N valori. Chiaramente si potranno formare almeno $N \times N \times \cdots N$ parole (n fattori). D'altra parte questo numero è altrettanto chiaramente il massimo consentito. Dunque si potranno formare N^n parole; se $n = N$ le parole saranno n^n .

Il caso delle successioni composte da elementi che valgono 0 o 1 corrisponde ad un alfabeto di due lettere, e infatti in questo caso le parole sono 2^n . Si noti che con la definizione precedente due parole coincidono se sono composte dalle stesse lettere dell'alfabeto *elencate nel medesimo ordine*. In altre parole basta scambiare la posizione di due lettere per avere due parole diverse. Inoltre non si tiene conto della semantica (cioè del fatto che la parola formata abbia un significato o no). Ad esempio, quindi, con l'alfabeto italiano (21 lettere, $N = 21$) si possono formare 21 parole con 1 lettera, $21^2 = 441$ parole con due lettere, $21^3 = 9261$ parole di tre lettere, $21^4 = 194481$ parole di 4 lettere, ecc. ; già nel caso delle parole composte da una sola lettera, però, osserviamo che quelle che hanno un significato sono solo 4 sulle 21 possibili (le 4 vocali a, e, i, o).

Esempio 5.13 (*Successioni aleatorie: il lancio della moneta*)

Supponiamo di lanciare una moneta, scrivendo 0 se viene testa e 1 se viene croce, e ripetiamo il lancio n volte. Otterremo una stringa di n elementi ciascuno dei quali vale 0 o 1: $\omega(n) = (0, 1, 1, \dots)$ dove la posizione di ogni 0 e ogni 1 è casuale e del tutto indipendente dalla

⁹Una delle figure più importanti nell'ideazione dei calcolatori elettronici e nell'elaborazione dei loro linguaggi fu il matematico John Von Neumann (Budapest 1903-Washington 1957).

posizione di ogni altro elemento precedente e successivo (in altre parole, ogni nuovo lancio "scorda" completamente il passato, cioè i risultati di tutti i lanci precedenti). Allo stesso modo, supponiamo di lanciare un dado. Qui scriveremo naturalmente 1 se esce 1, 2 se esce 2, ecc. Dopo n lanci, avremo una stringa $\omega(n) = (y_1, \dots, y_n)$ di lunghezza n dove $y_i \in \{1, 2, 3, 4, 5, 6\} \forall i$. Ad esempio, $\omega(n) = (3, 3, 1, 5, 5, 1, 1, 6, 4, \dots, 2)$.

Ancora più in generale, supponiamo di effettuare un esperimento qualsiasi che possa dare solo N risultati mutuamente esclusivi $x_1, \dots, x_N$ (cioè se viene x_k non può venire alcun altro x_j con $j \neq k$), detti *eventi elementari* e di ripetere n volte questo esperimento. Esempi: lanciamo 20 volte un dado: $n = 20$, $N = 6$ eventi elementari; lanciamo 30 volte la pallina della roulette: $n = 30$, $N = 36$ eventi elementari; misuriamo 100 volte il passaggio di una particella elementare in un rivelatore: $n = 100$, $N = 2$ eventi elementari (1 se passa, 0 se non passa). La successione delle estrazioni del lotto in un anno: $n = 52$, $N = 90$ eventi elementari.

Tutti gli *eventi* si possono definire a partire dagli eventi elementari. Esempi: lanciamo un dado. L'uscita di un numero pari è un evento definito dalla richiesta che il lancio dia come risultato uno qualsiasi dei risultati elementari rappresentati dai numeri 2, 4, 6. Allo stesso modo, un evento è l'uscita del rosso o del nero alla roulette; eventi sono l'estrazione di un numero qualsiasi fra 1 e 30 o di un numero qualsiasi fra 31 e 90 al lotto, oppure l'estrazione di un ambo, di una terna, quaterna, cinquina ecc. Un evento è il passaggio di una particella elementare in un rivelatore, ecc.

Poiché nella pratica le prove si ripetono tante volte, ci interessa ovviamente sapere cosa succede quando lo si fa il maggior numero di volte possibile, e soprattutto ci interessa *potere prevedere a priori* il risultato; matematicamente, vogliamo studiare cosa succede alla stringa quando $n \rightarrow \infty$.

Nel caso più semplice, quello della moneta, essa diventa una successione di elementi 0 e 1, ed è chiaro che non possiamo aspettarci che sia una successione convergente. Il limite infatti potrebbe valere solo 0 o 1. Se ne concluderebbe che dopo un certo numero *finito* di prove il lancio darebbe o definitivamente testa o definitivamente croce, il che è molto difficile da immaginare anche se la moneta è truccata a piacere.

Per studiare un pò più in dettaglio quello che succede conviene anzitutto introdurre due nozioni importanti: la *frequenza*, o più precisamente la *frequenza relativa* di avvenimento di

un evento qualsiasi Ω e la *probabilità* del verificarsi di un evento parimenti qualsiasi Ω :

Definizione 5.8

1. Dicesi *frequenza (relativa) dell'evento Ω* , denotata $\nu(\Omega)$, il numero di volte r in cui l'esperimento dà il risultato Ω (detto anche *numero degli esiti Ω*) diviso il numero totale n di ripetizioni dell'esperimento medesimo. In formule:

$$\nu(\Omega) = \frac{r}{n} \quad (5.9)$$

2. Dicesi *probabilità $p(\Omega)$ dell'evento Ω* il rapporto fra il numero dei casi favorevoli (cioè quelli in cui l'evento può verificarsi) f e il numero dei casi possibili N . In formule:

$$p(\Omega) = \frac{f}{N} \quad (5.10)$$

Osservazione 5.15

1. Ad esempio, nel caso del lancio di una moneta (non truccata) $N = 2$ e $f_k = 1$ tanto per la testa 0 quanto per la croce 1. Pertanto $p_0 = p_1 = 1/2$. Nel caso del dado (non truccato) $N = 6$ e $f_k = 1$ per ogni faccia. Quindi $p_1 = \dots = p_6 = 1/6$. Nel caso dell'evento che esca un numero pari nell'estrazione di un dado avremo $N = 6$, $f = 3$ e quindi $p = \frac{1}{2}$. Nel caso che la prima estrazione del lotto sia un numero fra dell'evento che l'estrazione del lotto dia un numero fra 1 e 30 avremo $f = 30$, $N = 90$ e quindi $p_1 = \frac{1}{3}$, mentre sarà $p_2 = \frac{60}{90} = \frac{2}{3}$ la probabilità che venga estratto un numero fra 31 e 90.
2. La probabilità di un evento qualsiasi è per sua natura un numero compreso fra 0 (in corrispondenza all'evento impossibile, cioè l'evento che non può mai avvenire: ad esempio, che lanciando un dado non venga nessuno dei numeri 1, 2, 3, 4, 5, 6) e 1 (in corrispondenza all'evento certo, cioè l'evento che avviene sicuramente: ad esempio, nel lancio del dato, che venga uno qualsiasi dei numeri 1, 2, 3, 4, 5, 6).
3. Siano dati due eventi, uno con probabilità p_1 e un altro con probabilità p_2 . Supponiamo che essi si escludano a vicenda, cioè se si verifica l'uno non può verificarsi l'altro e viceversa. Allora la probabilità p che si verifichi *uno qualsiasi dei due* vale la somma

- delle probabilità che si verifichi o l'uno o l'altro: $p = p_1 + p_2$. Ad esempio, la probabilità che ad un giro di roulette un numero compreso fra 7 e 12 oppure fra 19 e 30 è $\frac{6}{36} + \frac{12}{36} = \frac{1}{2}$.
4. Se possono avvenire solo un numero finito N di eventi x_k che si escludono a vicenda, e ciascuno di questi ha probabilità p_k , dovrà essere $\sum_{k=1}^N p_k = 1$ perché almeno uno degli eventi si verifica sicuramente.
5. Siano dati ancora due eventi, uno con probabilità p_1 e un altro con probabilità p_2 . Supponiamo che siano indipendenti, cioè che il verificarsi dell'uno non abbia alcuna influenza sul verificarsi dell'altro e viceversa. Allora la probabilità che essi si verifichino *simultaneamente* vale il prodotto delle probabilità che si verifichi o l'uno o l'altro: $p = p_1 \cdot p_2$. Ad esempio, la probabilità che ad un giro di roulette venga estratto un numero rosso fra 20 e 31 vale $\frac{1}{2} \times \frac{12}{36} = \frac{1}{6}$. La probabilità che gettando un dado venga un numero pari non superiore a 4 vale $\frac{4}{6} \cdot \frac{1}{2} = \frac{1}{3}$. Poiché nelle estrazioni del lotto si scartano dalle estrazioni successive i numeri che di volta in volta escono, la probabilità di fare un ambo giocando al lotto su una sola ruota è evidentemente $\frac{1}{90} \times \frac{1}{89} = \frac{1}{8010}$, quella di fare un terno $\frac{1}{90} \times \frac{1}{89} \times \frac{1}{88} = \frac{1}{704880}$, ecc.
6. Mentre la frequenza è una quantità che si può determinare solo a posteriori, cioè dopo avere fatto gli esperimenti, la probabilità è una quantità che si definisce a priori.

La *legge dei grandi numeri* (che qui riportiamo senza dimostrazione nella formulazione adattata al caso particolare qui considerato) consiste in sostanza nell'affermazione (molto intuitiva) che ripetendo le prove un numero di volte sempre crescente la frequenza di avvenimento di ogni singolo evento tende alla probabilità a priori che esso avvenga, *purché le prove siano del tutto indipendenti l'una dall'altra*:

Dato $\delta > 0$, sia P_n la probabilità che nella stringa $\omega(n)$ risulti $|\nu(\Omega) - p(\Omega)| > \delta$ per almeno un evento Ω . Allora la successione P_n tende a 0 per $n \rightarrow \infty$.

Osservazione 5.16

1. Poiché la probabilità è definita a priori, la legge dei grandi numeri ci permette in un certo senso di *prevedere* i risultati di un esperimento ripetuto un gran numero di volte

in maniera indipendente. Una delle applicazioni più importanti è questa: se dopo un gran numero di prove la frequenza si discosta dalla probabilità, c'è fondata ragione di pensare che esse non siano indipendenti; nel linguaggio della fisica, nel compiere l'esperimento si sta commettendo un *errore sistematico*.

2. Bisogna però guardarsi bene dal trarre conseguenze avventate dalla legge dei grandi numeri. Una delle più comuni è quella di pensare che la n -esima estrazione "ricordi" le precedenti, nel senso che se lanciando la moneta ad esempio 10 volte è venuto 10 volte testa nel lancio successivo l'evento croce sia molto più probabile dell'evento testa. La probabilità rimane invece la stessa (ovviamente se la moneta non è truccata). Questa errata interpretazione della legge dei grandi numeri sta alla base di molte perdite al gioco.

Questo concetto di convergenza "in probabilità" può apparire oscuro. Conviene riformularlo in termini più vicini a quelli fin qui esposti, introducendo una nozione di convergenza assai utile, quella delle medie aritmetiche isolata da Cesàro¹⁰ più di un secolo fa.

Definizione 5.9

1. Dati n numeri reali qualsiasi $r_1, \dots, r_n$, dicesi loro media aritmetica la quantità:

$$M_n = \frac{1}{n} \sum_{k=1}^n r_k \quad (5.11)$$

2. Data la successione a_k , dicesi successione delle medie aritmetiche da essa dedotte la successione

$$M_1 = a_1, \quad M_2 = \frac{a_1 + a_2}{2}, \quad M_3 = \frac{a_1 + a_2 + a_3}{3}, \quad \dots, \quad M_n = \frac{1}{n} \sum_{k=1}^n a_k, \dots$$

Definizione 5.10 La successione a_k converge a l nel senso di Cesàro se la successione M_n delle medie aritmetiche da essa dedotte converge a l , cioè se

$$\lim_{n \rightarrow \infty} M_n = \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=1}^n a_k = l$$

¹⁰Ernesto Cesàro (Napoli 1859-Torre Annunziata 1906), Professore di Analisi matematica all'Università di Napoli.

Osservazione 5.17

La nozione di convergenza come è stata definita sopra (detta anche convergenza *semplice*) è *più forte* della nozione di convergenza nel senso di Cesàro. Equivalentemente, la nozione di convergenza secondo Cesàro è *più debole* di quella abituale. Questo significa che se una successione converge nel senso abituale converge anche nel senso di Cesàro, mentre *non vale il viceversa*. Infatti:

Proposizione 5.1

1. Se una successione a_n converge, $\lim_{n \rightarrow \infty} a_n = l$ allora anche $\lim_{n \rightarrow \infty} M_n = l$.
2. Esistono successioni limitate a_n tali che $\lim_{n \rightarrow \infty} M_n = l \in \mathbb{R}$ ma che non ammettono limite.

Dimostrazione.

1. È sufficiente dimostrare l'asserzione per $l = 0$. Altrimenti basta considerare la successione $a_n - l$ la cui successione delle medie aritmetiche è $M_n - l$: infatti se $a_n - l \rightarrow 0$ ciò implica $M_n - l \rightarrow 0$ e quindi $M_n \rightarrow l$. Dato $\epsilon > 0$, poiché a_n è convergente esiste $\nu(\epsilon)$ tale che $|a_n| < \frac{\epsilon}{2}$ se $n > \nu(\epsilon)$. Allora, usando la disuguaglianza triangolare $|x + y| \leq |x| + |y|$ valida $\forall x, y \in \mathbb{R}$:

$$\left| \frac{1}{n} \sum_{k=1}^n a_k \right| \leq \left| \frac{1}{n} \sum_{k=1}^{\nu} a_k \right| + \left| \frac{1}{n} \sum_{k=\nu+1}^n a_k \right| < \left| \frac{1}{n} \sum_{k=1}^{\nu} a_k \right| + \frac{n - \nu}{n} \frac{\epsilon}{2} < \left| \frac{1}{n} \sum_{k=1}^{\nu} a_k \right| + \frac{\epsilon}{2}$$

perché $\frac{n - \nu}{n} < 1$. Sia ora $M := \max |a_k| : k = 1, \dots, \nu$. Allora

$$\left| \frac{1}{n} \sum_{k=1}^{\nu} a_k \right| \leq \frac{M\nu}{n} + \frac{\epsilon}{2}$$

Ora il primo addendo non supererà $\frac{\epsilon}{2}$ se $n > \frac{2M\nu}{\epsilon}$. Sia ora $n_1(\nu) := \max \left(\nu(\epsilon), \frac{2M\nu}{\epsilon} \right)$.

Allora $\forall n > n_1(\nu)$ si avrà:

$$\left| \frac{1}{n} \sum_{k=1}^n a_k \right| \leq \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon$$

e ciò dimostra la convergenza a 0 della successione M_n .

2. Abbiamo già visto che la successione $0, 1, 0, 1, \dots$ non converge. Facciamo vedere che invece converge a $\frac{1}{2}$ nel senso di Cesàro. Infatti si ha:

$$M_{2p} = \frac{1}{2p} \sum_{k=1}^{2p} a_k = \frac{p}{2p} = \frac{1}{2}, \quad p = 1, 2, \dots \quad M_{2p+1} = \frac{1}{2p+1} \sum_{k=1}^{2p+1} a_k = \frac{p}{2p+1} \quad p = 1, 2, \dots$$

Ora $\lim_{p \rightarrow \infty} M_{2p+1} = \lim_{p \rightarrow \infty} \frac{1}{2 + 1/p} = \frac{1}{2}$ perché $\lim_{p \rightarrow \infty} (2 + \frac{1}{p}) = 2 + \lim_{p \rightarrow \infty} \frac{1}{p} = 2$ (il limite della somma vale la somma dei limiti), e $\lim_{p \rightarrow \infty} \frac{1}{2 + 1/p} = \frac{1}{\lim_{p \rightarrow \infty} (2 + 1/p)}$ (il limite del quoziente vale il quoziente dei limiti se il limite dei denominatori non è 0). Dunque: $\lim_{n \rightarrow \infty} M_n = \frac{1}{2}$.

Abbiamo così trovato un esempio di successione convergente nel senso di Cesàro ma non in senso semplice e ciò dimostra la proposizione.

Esercizio 5.18

Trovare il limite secondo Cesàro della successione $1, -1, 1, -1, \dots$ cioè $a_n = 1$ se $n = 2p$ e $a_n = -1$ se $n = 2p + 1$, $p = 0, 1, \dots$

Soluzione. 0.

Vogliamo ora applicare la nozione di convergenza nel senso di Cesàro all'esame del limite per $n \rightarrow \infty$ delle nostre stringhe $\omega(n) = (y_1, y_2, \dots, y_n)$ dove ogni y_k può assumere un numero finito di valori $x_1, \dots, x_N$; ciascuno di questi può avvenire con probabilità p_k . Restringiamo ora la nostra attenzione ad un particolare evento elementare. Senza ledere la generalità possiamo assumere che sia y_1 , perché per tutti gli altri si ragiona esattamente allo stesso modo. A partire dalla stringa $\omega(n) = (y_1, y_2, \dots, y_n)$ di lunghezza n definiamo una stringa $\tau(n)$ di pari lunghezza nel modo seguente: tutte le volte che uno degli y nella stringa prende il valore x_1 scriviamo 1, altrimenti scriviamo 0. In altre parole la stringa $\tau(n)$ è una stringa di n elementi 0 o 1 definita a partire dagli esiti delle prove ripetute scrivendo 1 se l'esito è x_1 , e scrivendo 0 se non lo è.

In formule, se $\tau_h(n)$ denota l'elemento di posto h nella stringa $\tau(n)$, si ha

$$\tau_h^1(n) = \begin{cases} 1 & \text{se } y_h = x_1 \\ 0 & \text{se } y_h \neq x_1 \end{cases}$$

ovvero, usando il simbolo di Kronecker ¹¹ $\delta_{x,y}$: $\tau_h^1(n) = \delta_{y_h, x_1}$. Allo stesso modo si definiscono le altre stringhe $\tau_h^k(n)$, $k = 2, \dots, N$. Ad esempio, consideriamo i lanci ripetuti di un dado, $N = 6$. Gli eventi elementari sono l'uscita dei numeri 1, 2, 3, 4, 5, 6. Concentriamoci ad esempio sull'uscita del numero 2, scrivendo 1 se viene e 0 se non viene. Otterremo allora una successione del tipo 000100110.... La stessa cosa vale per gli altri numeri diversi da 2. È

¹¹ $\delta_{x,y} = 1$ se $x = y$; $\delta_{x,y} = 0$ se $x \neq y$.

chiaro che ci sarà da aspettarsi che, se la stringa è abbastanza lunga, in media la proporzione fra gli 1 e gli 0 sarà di 1 a 6.

Infatti la legge dei grandi numeri può essere enunciata anche in questo modo:

Dato $\delta > 0$, sia P_n la probabilità che nella stringa $\tau^k(n)$ risulti $|M_n^k(\tau) - p_k| > \delta$ per almeno un k . Allora la successione P_n tende a 0 per $n \rightarrow \infty$.

$M_n^k(\tau)$ è ovviamente la successione delle medie aritmetiche dedotte da $\tau^k(n)$.

Per dimostrare questa formulazione, ammettendo quella precedente, è sufficiente osservare che $M_n(\tau) = \nu(x_k)$. Infatti:

$$M_n(\tau) = \frac{1}{n} \sum_{h=0}^n \tau_h(n) = \frac{f_k}{n} = \nu(x_k)$$

dove f_k è il numero di volte in cui l'esito è x_k . Pertanto, possiamo enunciare la legge dei grandi numeri anche così:

La probabilità P_n che la media di Cesàro della stringa $\tau^k(n)$ differisca dalla probabilità p_k per una quantità piccola a piacere tende a 1 al tendere all'infinito del numero delle prove n .

In altre parole, la probabilità che una successione aleatoria di elementi 0 e 1 del tipo definito in precedenza abbia limite nel senso di Cesàro tende alla certezza al tendere all'infinito del numero delle prove, purché indipendenti.

6 Frazioni continue

Un'altra realizzazione concreta molto importante del concetto di successione è la nozione di frazione continua.

Il modo migliore per introdurla è richiamare l'algoritmo di Euclide¹² per calcolare il massimo comune divisore (MCD) fra due numeri naturali.

¹²Euclide, matematico greco vissuto ad Alessandria nella prima metà del III secolo avanti Cristo. Di lui si sa molto poco, se non che insegnava matematica al Museo di Alessandria, la più importante istituzione scientifica dell'epoca ellenistica. Esistono numerose e valide ragioni indirette per ritenere che Archimede ed Eratostene abbiano studiato sotto suoi allievi. Fu autore degli *Elementi* e dell'*Ottica*. Gli *Elementi* di Euclide rappresentano il primo vero libro di matematica mai scritto, e a parere di molti il più importante di tutti i tempi. Furono il libro di testo fondamentale per la matematica già nell'antichità. Conservati dalla tradizione bizantino-araba dopo la caduta dell'impero romano d'occidente, riemersero in Occidente col Rinascimento e ripresero il loro posto di libro di testo fondamentale nelle scuole di molti paesi, inclusa l'Italia, fino quasi alla fine dell'ottocento.

Esempio 6.14 (*Algoritmo di Euclide per calcolare il massimo comune divisore fra due numeri naturali*)

Facciamo prima un paio di esempi numerici.

1. Calcolare il massimo comune divisore (MCD) fra 13 e 8. L'algoritmo di Euclide procede così:

$$8 \text{ sta in } 13 \text{ una volta con resto di } 5, \text{ cioè } \frac{13}{8} = 1 + \frac{5}{8}.$$

$$5 \text{ sta in } 8 \text{ una volta con resto di } 3: 8 = 1 \cdot 5 + 3, \text{ cioè } \frac{8}{5} = 1 + \frac{3}{5}.$$

$$3 \text{ sta una volta in } 5 \text{ con resto di } 2: 5 = 1 \cdot 3 + 2, \text{ cioè } \frac{5}{3} = 1 + \frac{2}{3}.$$

$$2 \text{ sta in } 3 \text{ una volta con resto di } 1, 3 = 1 \cdot 2 + 1, \text{ cioè } \frac{3}{2} = 1 + \frac{1}{2}.$$

$$1 \text{ sta in } 2 \text{ esattamente } 2 \text{ volte e l'algoritmo si arresta: } 2 = 2 \cdot 1 + \frac{0}{2}.$$

In parole: si divide il maggiore per il minore tenendo il resto; poi ad ogni passo si divide il divisore per il resto, tenendo il secondo resto finché non si trova un quoziente esatto (resto 0).

Riassumendo:

$$\frac{13}{8} = 1 + \frac{5}{8}; \quad \frac{8}{5} = 1 + \frac{3}{5}, \quad \frac{5}{3} = 1 + \frac{2}{3} \quad \frac{3}{2} = 1 + \frac{1}{2} \quad \frac{2}{1} = 2 \cdot 1 + 0 \quad (6.12)$$

Il numero a cui l'algoritmo si arresta (cioè arriva al resto 0) è il MCD, qui 1, come sapevamo già perché 13 è primo.

2. Calcoliamo come secondo esempio il MCD fra 36 e 15. Si ha

$$36 = 2 \cdot 15 + 6, \quad 15 = 2 \cdot 6 + 3, \quad 6 = 2 \cdot 3$$

ovvero

$$\frac{36}{15} = 2 + \frac{6}{15}; \quad \frac{15}{6} = 2 + \frac{3}{6}; \quad \frac{6}{3} = 2 + 0$$

e quindi il MCD è 3.

La formulazione generale dell'algoritmo di Euclide sarà quindi la seguente:

Dati due numeri naturali $p > q$, per calcolare il loro massimo comun divisore M si procede induttivamente a partire da $p : q$ tramite divisioni successive del divisore per il resto:

$$p = s_0q + r_1$$

$$q = s_1r_1 + r_2$$

$$r_1 = s_2r_2 + r_3$$

$$r_2 = s_3r_3 + r_4$$

.....

ovvero

$$\frac{p}{q} = s_0 + \frac{r_1}{q}; \quad \frac{q}{r_1} = s_1 + \frac{r_2}{r_1}; \quad \frac{r_1}{r_2} = s_2 + \frac{r_3}{r_2} \dots$$

Poiché p e q sono finiti, esiste un passo n al quale il procedimento si arresta nel senso che la divisione dà un risultato esatto: $r_{n-1} = s_n r_n + 0$ con $r_n \geq 1$. In tal caso r_n è il MCD. p e q sono primi fra loro se e solo se $r_n = 1$.

Osservazione 6.18

Il modo consueto di calcolare il MCD fra p e q è tramite la scomposizione in fattori primi: il MCD è il prodotto di *tutti* i fattori primi comuni (tutti nel senso che ciascuno va contato con la sua molteplicità). Ad esempio, il MCD fra $600 = 2^3 \cdot 3 \cdot 5^2$ e $490 = 2 \cdot 5 \cdot 7^2$ è $2 \cdot 5 = 10$. Il MCD fra $11520 = 2^7 \cdot 3^2 \cdot 5$ e $120 = 2^3 \cdot 3 \cdot 5$ è $2^3 \cdot 3 = 24$.

L'algoritmo di Euclide, apparentemente più complicato, è però molto più rapido perché è un metodo iterativo e non enumerativo come la scomposizione in fattori primi. Dato un numero naturale qualsiasi, infatti, la sua scomposizione in fattori primi richiede l'esame della divisibilità per 2, poi per 3, poi per 5, ecc. In particolare l'algoritmo di Euclide è molto più efficiente sul calcolatore, perché si può dimostrare che, dati due numeri naturali ($n > m$), che il tempo impiegato dal calcolatore per calcolarne il MCD col metodo enumerativo cresce esponenzialmente con n , mentre quello impiegato dall'algoritmo di Euclide cresce come $\ln n$.

Esercizio 6.19

Calcolare il MCD fra le seguenti coppie di numeri

$$(8271, 3215); \quad (6632, 2783); \quad (9215, 1003); \quad (1200, 6557)$$

(a) Con l'algoritmo di Euclide.

(b) Con l'usuale scomposizione in fattori primi.

In ciascuno dei casi, con quale metodo vi sembra di avere fatto prima?

Osservazione 6.19

L'algoritmo di Euclide può essere applicato senza alcuna variazione per trovare il MCD fra due polinomi. Basta solo sostituire la divisione fra polinomi alla divisione fra numeri interi. Per i dettagli, si segua il corso di Algebra.

Consideriamo ora la frazione $\frac{p}{q}$. Usando le divisioni con resto definite dall'algoritmo di Euclide possiamo scrivere:

$$\begin{aligned} \frac{p}{q} &= s_0 + \frac{r_1}{q} = s_0 + \frac{1}{s_1 + \frac{r_2}{r_1}} = s_0 + \frac{1}{s_1 + \frac{1}{s_2 + \frac{r_3}{r_2}}} \\ &= \dots = s_0 + \frac{1}{s_1 + \frac{1}{s_2 + \frac{1}{s_3 + \frac{1}{\dots + \frac{1}{s_n}}}}} \end{aligned} \quad (6.13)$$

Ad esempio, dalla (6.12) ricaviamo $s_0 = 1, s_1 = 1, s_2 = 1, s_3 = 1, s_4 = 2$ per $\frac{13}{8}$, e $s_0 = 2, s_1 = 2, s_3 = 2$ per $\frac{36}{15}$. Infatti si verificano subito le uguaglianze

$$1 + \frac{1}{1 + \frac{1}{1 + \frac{1}{1 + \frac{1}{2}}}} = \frac{13}{8}; \quad 2 + \frac{1}{2 + \frac{1}{2}} = \frac{12}{5} = \frac{36}{15}$$

Definizione 6.11 *L'espressione (6.13) si dice sviluppo in frazione continua aritmetica (finita) del numero razionale $\frac{p}{q}$, e ne rappresenta la frazione ridotta ai minimi termini; si usa anche la notazione abbreviata: $\frac{p}{q} = [s_0, s_1, \dots, s_n]$.*

(Ad esempio, $\frac{13}{8} = [1, 1, 1, 1, 2]$; $\frac{36}{15} = [2, 2, 2]$).

Esercizio 6.20

Si ricavino i seguenti sviluppi in frazione continua:

$$\frac{56}{24} = [2, 3] = 2 + \frac{1}{3}; \quad \frac{120}{75} = [1, 1, 1, 2] = 1 + \frac{1}{1 + \frac{1}{1 + \frac{1}{2}}}; \quad \frac{125}{64} = [1, 1, 20, 3] = 1 + \frac{1}{1 + \frac{1}{20 + \frac{1}{3}}}$$

Un esempio molto interessante di impiego delle frazioni continue finite è il calcolo del calendario di uso quotidiano.

Esempio 6.15 (*Il calendario*)

Si dice *anno comune* la durata di 365 giorni esatti. Si dice *anno solare* la durata esatta dell'intervallo di tempo che la terra impiega a compiere una rivoluzione completa attorno al sole. Ebbene l'anno solare dura 365 giorni, 5 ore, 48 minuti e 49 secondi, cioè 5 ore, 48 minuti e 49 secondi più dell'anno comune. Poiché per evidenti necessità pratiche si vuole elaborare il calendario in modo che l'anno contenga un numero che differisca il meno possibile da 365, e d'altra parte si deve per forza tenere conto di questa maggiore durata, occorrerà *intercalare* ogni tanto dei giorni. Intercalare significa aggiungere un giorno ogni qualche anno, cioè fare in modo che, a scadenze fisse da calcolare e da condensare in una regola molto semplice da capire e da applicare, l'anno duri 366 giorni invece di 365 (anno bisestile).

Come ottenere questo risultato con la massima precisione? Riducendo tutto a secondi, possiamo scrivere

$$\frac{24 \text{ ore} = 1 \text{ giorno}}{5h 48'49''} = \frac{86400}{20929}$$

Ne deduciamo che intercalando 20929 giorni nell'arco di 86400 anni comuni otteniamo precisamente 86400 anni solari, e questa è la minima durata temporale che risulti pari ad un multiplo intero dell'anno solare.

Sviluppiamo ora in frazione continua il numero razionale $\frac{86400}{20929}$. Si ha:

$$\begin{aligned} \frac{86400}{20929} &= [4, 7, 1, 3, 1, 16, 1, 1, 15] \\ &= 4 + \frac{1}{7 + \frac{1}{1 + \frac{1}{3 + \frac{1}{1 + \frac{1}{16 + \frac{1}{1 + \frac{1}{1 + \frac{1}{15}}}}}}}} \end{aligned}$$

Calcoliamo la successione delle frazioni parziali: essa è

$$\begin{aligned} 4 & \quad ; \quad \frac{29}{7} = 4,1428571429; \quad \frac{33}{8} = 4,125; \quad \frac{128}{31} = 4,1290322521; \\ \frac{161}{39} & = 4,1282051282; \quad \frac{2704}{655} = 4,1282442748; \quad \frac{2865}{694} = 4,1282420749; \\ \frac{5569}{1349} & = 4,1282431431; \quad \frac{86400}{20929} = 4,1282431076 \end{aligned}$$

Dunque la prima e più semplice intercalazione è l'aggiunta di un giorno ogni 4 anni. Questo è il calendario giuliano¹³ secondo il quale un anno ogni 4 è bisestile. Un'intercalazione migliore sarebbe quella di 7 giorni da inserire nell'arco di 29 anni, e una ancora migliore sarebbe quella di 8 giorni da inserire nell'arco di 33 anni. Quest'ultima è l'intercalazione dell'anno persiano. Il calendario detto gregoriano¹⁴ in uso oggi in tutto il mondo fa intervenire tuttavia un'intercalazione diversa, e precisamente 97 giorni inseriti nell'arco di 400 anni secondo la regola che ricorderemo fra poco. Poiché $\frac{400}{97} = 4,1237113402$ essa è più accurata di quella di 7 giorni nell'arco di 29 anni ma *meno* accurata di quella di 8 giorni nell'arco di 33 anni del calendario persiano. I motivi sono due. Il primo è che i dati astronomici a disposizione di Clavio sulla durata dell'anno solare erano meno precisi di quelli degli astronomi ellenistici sui quali si erano basati i persiani. Il secondo e più importante motivo consiste nel fatto che l'intercalazione gregoriana è condensabile in una regola più semplice di quella del calendario persiano perché fa riferimento solo ai secoli. Basta infatti, come tutti sappiamo, apportare la seguente semplicissima modifica al calendario giuliano: considerare bisestili gli anni secolari solo se le loro prime due cifre sono un numero divisibile per 4. Ad esempio il 1700, il 1800 e il 1900 non sono stati bisestili, mentre il 1600 e il 2000 lo sono stati. Nel calendario persiano invece si considerano otto cicli in cui un anno ogni 4 è bisestile, poi un ciclo di cinque anni

¹³Fatto elaborare agli astronomi alessandrini da Giulio Cesare in occasione della sua spedizione in Egitto. È ancora in vigore per la Chiesa Ortodossa.

¹⁴Dal nome del Papa Gregorio XIII (Ugo Boncompagni), bolognese, che regnò dal 1573 al 1585. La sua statua bronzea sta sulla facciata del Palazzo d'Accursio in Piazza Maggiore a Bologna. Egli lo promulgò nel 1583 dichiarando inesistenti i giorni dal 4 al 15 ottobre. Un ritardo di meno di 12 minuti all'anno può sembrare trascurabile, ma diventa apprezzabile al trascorrere dei secoli. Al momento della promulgazione del calendario gregoriano in vigore oggi si era accumulato appunto un ritardo di 12 giorni che scombinava le date degli equinozi; ne veniva alterato l'intero calendario ecclesiastico, che coincideva a quei tempi con quello civile. Per questo motivo il papa decise di riformarlo. Il calendario gregoriano fu elaborato dal gesuita Cristoforo Clavio (italianizzazione di Christophorus Clavius, a sua volta latinizzazione di Christoph Klau, Bamberg 1538-Roma 1612), matematico capo del Collegio Romano. Gli appassionati di cinema ricorderanno che la stazione orbitante fra la Terra e la Luna in *2001, Odissea nello spazio* si chiamava Clavius. Il suo principale collaboratore in quest'opera fu un altro gesuita, Gian Antonio Maggini, Professore poi all'Università di Bologna dal 1592 (anno in cui fu preferito all'allora giovanissimo Galileo).

di cui l'ultimo è bisestile, per poi ricominciare con gli otto cicli di 4 anni.

È evidente che una frazione continua finita definisce un numero razionale ed uno solo; d'altra parte ogni numero razionale ammette un unico sviluppo in frazione continua finita per l'algoritmo di Euclide. Dunque frazioni continue aritmetiche finite e numeri razionali sono in corrispondenza biunivoca. Infatti le prime dimostrazioni dell'irrazionalità di e (Eulero, 1737) e di π (Lambert, 1772) consistevano nel far vedere che essi ammettevano sviluppi in frazione continua *infinita*.

Vediamo ora come generalizzando l'algoritmo di Euclide ai numeri reali si ottengono gli sviluppi in frazione continua infinita.

Se x è un numero non intero, si continua a denotare $[x]$ denota la *parte intera* di x , e $\{x\}$ la *parte decimale*:

$$x = [x] + \{x\}.$$

Ovviamente $[x] \in \mathbb{Z}$, e $0 \leq \{x\} < 1$. Allora possiamo scrivere $\{x\} = \frac{1}{x_1}$ con $x_1 > 1$.

Denotando $[x] = a_0$ abbiamo $x = a_0 + \frac{1}{x_1}$. Ora, allo stesso modo, possiamo scrivere

$$x_1 = a_1 + \frac{1}{x_2}, \quad x_2 = a_2 + \frac{1}{x_3}, \dots \quad (6.14)$$

ossia

$$x = a_0 + \frac{1}{a_1 + \frac{1}{x_2}} = [a_0, a_1, x_2]; \quad x = a_0 + \frac{1}{a_1 + \frac{1}{a_2 + \frac{1}{x_3}}} = [a_0, a_1, a_2, x_3]$$

Iterando n volte il procedimento si avrà

$$x = a_0 + \frac{1}{a_1 + \frac{1}{a_2 + \frac{1}{a_3 + \frac{1}{\dots + \frac{1}{x_n}}}}} = [a_0, a_1, \dots, x_n] \quad (6.15)$$

In generale:

$$x_n = [x_n] + \{x_n\} = a_n + \frac{1}{x_{n+1}}$$

dove $a_n = [x_n]$.

Se $x \notin \mathbb{Q}$ sappiamo, per l'osservazione precedente, che non esiste alcun $n \in \mathbb{N}$ per cui il procedimento si arresta. Ciò significa che non esiste alcun n tale che $\{x_n\} = 0$, cosicché si può sempre porre $\{x_n\} = \frac{1}{x_{n+1}}$. Pertanto la successione $[a_0, a_1, \dots, x_n]$ può essere continuata indefinitamente. Il procedimento di Euclide definisce quindi una successione $[a_0, a_1, \dots, a_n, \dots]$ detta *frazione continua infinita aritmetica*, o semplicemente *frazione continua aritmetica* generata da x , che si indica col simbolo:

$$a_0 + \frac{1}{a_1 + \frac{1}{a_2 + \frac{1}{a_3 + \frac{1}{\dots + \frac{1}{a_n + \frac{1}{\dots}}}}}} \quad (6.16)$$

Più in generale, date due successioni a_k, b_k di numeri reali, per semplicità positivi, consideriamo la successione J_n definita induttivamente nel modo seguente:

$$\begin{aligned} J_1 &= \frac{b_1}{a_1}; & J_2 &= \frac{b_1}{a_1 + \frac{b_2}{a_2}} & J_3 &= \frac{b_1}{a_1 + \frac{b_2}{a_2 + \frac{b_3}{a_3}}} \\ & & & & & \dots & J_n &= \frac{b_1}{a_1 + \frac{b_2}{a_2 + \frac{b_3}{a_3 + \frac{b_4}{a_4 + \frac{b_5}{\dots + \frac{b_n}{a_n}}}}}} \end{aligned} \quad (6.17)$$

Definizione 6.12 *Dicesi frazione continua la successione J_n della (6.17). Se la successione converge a $J \in \mathbb{R}$, $\lim_{n \rightarrow \infty} J_n = J$, la frazione continua si dice convergente a J e si usa la notazione:*

$$J = \frac{b_1}{a_1 + \frac{b_2}{a_2 + \frac{b_3}{a_3 + \frac{b_4}{a_4 + \dots}}}} \quad (6.18)$$

Se tronchiamo la frazione continua a $n = 1, 2, \dots$ otteniamo per definizione la successione dei *troncamenti*, o degli *approssimanti*, o dei *convergenti*. La successione dei troncamenti coincide con la successione (6.17) e sarà ovviamente della forma $\frac{A_n}{B_n}$, dove i numeri positivi A_n

e B_n si dicono *numeratori parziali* e *denominatori parziali* degli approssimanti. La frazione continua si dice poi *aritmetica* se tutti i coefficienti a_k, b_k sono numeri naturali, e *periodica* se entrambe le successioni a_k e b_k sono periodiche. Se $b_k = 1 \forall k$ si scrive anche, come sopra, $J = [a_0, a_1, a_2, \dots]$.

Osservazione 6.20

Si può dimostrare che se $x \notin \mathbb{Q}$, allora $x = [a_0, a_1, a_2, \dots]$ dove $[a_0, a_1, a_2, \dots]$ è la frazione continua (6.16). In altre parole, la frazione continua generata da x converge a x .

Esempio 6.16

1. Calcoliamo la successione $[a_0, a_1, \dots, a_n, \dots]$, cioè la frazione continua del tipo (6.16) corrispondente a $x = \sqrt{2}$. Si ha:

$$\begin{aligned} \sqrt{2} &= 1 + \frac{1}{x_1} \Rightarrow x_1 = \frac{1}{\sqrt{2} - 1} = \sqrt{2} + 1; \\ x_1 &= 2 + \frac{1}{x_2} \Rightarrow x_2 = \frac{1}{x_1 - 2} = \frac{1}{\sqrt{2} - 1} = \sqrt{2} + 1 = x_1, x_3 = x_2, \dots \end{aligned}$$

Quindi la frazione continua è

$$1 + \frac{1}{2 + \frac{1}{2 + \frac{1}{2 + \frac{1}{2 + \dots}}}} \Rightarrow \sqrt{2} - 1 = [2, 2, 2, \dots] \quad (6.19)$$

2. Calcoliamo la successione $[a_0, a_1, \dots, a_n, \dots]$ corrispondente a $x = \frac{\sqrt{5} + 1}{2}$. Si ha

$$\frac{\sqrt{5} + 1}{2} = 1 + \frac{1}{x_1} \Rightarrow x_1 = \frac{2}{\sqrt{5} - 1} = \frac{\sqrt{5} + 1}{2}; \quad \frac{\sqrt{5} + 1}{2} = 1 + \frac{1}{x_2} \Rightarrow x_2 = x_1, \dots$$

Quindi la frazione continua è

$$\frac{\sqrt{5} + 1}{2} = \frac{1}{1 + \frac{1}{1 + \frac{1}{1 + \frac{1}{1 + \dots}}}} = [1, 1, 1, \dots] \quad (6.20)$$

Esercizio 6.21

Calcolare gli sviluppi in frazione continua del tipo $[a_0, a_1, a_2, \dots]$ di $\sqrt{3}$, $\sqrt{7}$, $\sqrt{11}$, $\sqrt{13}$.

Risposta.

$$\begin{aligned}\sqrt{3} &= [1, 1, 2, 1, 2, 1, 2, \dots]; & \sqrt{7} &= [2, 1, 1, 1, 4, 1, 1, 1, 4, 1, 1, 1, 4, 1, 1, 1, 4, \dots]; \\ \sqrt{11} &= [3, 3, 6, 3, 6, 3, \dots]; & \sqrt{13} &= [3, 1, 1, 1, 1, 1, 6, 1, 1, 1, 1, 1, 6, 1, \dots]\end{aligned}$$

Tutti gli esempi precedenti sono frazioni continue aritmetiche periodiche e convergenti per il teorema di Lagrange che ricorderemo fra poco. Esse rappresentano così lo sviluppo in frazione continua di $\sqrt{2}$ e $\frac{\sqrt{5}+1}{2}$, $\sqrt{3}$, $\sqrt{7}$, $\sqrt{11}$, $\sqrt{13}$ rispettivamente.

Il calcolo della successione dei convergenti di una frazione continua è molto più semplice di quanto potrebbe sembrare a prima vista. Basta infatti, nota la successione a_k dei coefficienti, calcolare solo i primi due numeratori e denominatori parziali per ottenere tutti gli altri per ricorrenza:

Esercizio 6.22

Dimostrare che i convergenti $\frac{A_n}{B_n}$ di ogni frazione continua aritmetica soddisfano la formula di ricorrenza

$$\frac{A_{n+1}}{B_{n+1}} = \frac{A_n a_n + A_{n-1}}{B_n a_n + B_{n-1}}$$

dove A_n e B_n sono successioni di numeri interi che a loro volta soddisfano le relazioni di ricorrenza

$$A_n := A_{n-1}a_{n-1} + A_{n-2}; \quad B_n := B_{n-1}a_{n-1} + B_{n-2}$$

con $A_0 = 1$, $A_1 = a_0 a_1$, $B_0 = 0$, $B_1 = a_1$.

Soluzione Ricordiamo che la frazione continua si ottiene dalla ricorrenza (6.14) a partire dalla condizione iniziale $x = a_0 + \frac{1}{x_1}$. Portando allo stesso denominatore otteniamo

$$x = a_0 + \frac{1}{x_1} \implies x = \frac{a_0 x_1 + 1}{x_1}, \quad x_1 = a_1 + \frac{1}{x_2} \implies x = \frac{(a_0 a_1 + 1)x_2 + 1}{a_1 x_2 + 1}, \dots$$

Dimostriamo allora per induzione che si ha

$$x = \frac{P_n x_n + P_{n-1}}{Q_n x_n + Q_{n-1}}$$

dove P_n e Q_n sono successioni di numeri interi. L'affermazione è vera per $n = 1$; assumiamola vera per $n - 1$. Allora le formule

$$x = \frac{P_{n-1}x_{n-1} + P_{n-2}}{Q_{n-1}x_{n-1} + Q_{n-2}}; \quad x_{n-1} = a_{n-1} + \frac{1}{x_n}$$

implicano

$$x = \frac{P_{n-1}a_{n-1} + \frac{P_{n-1}}{x_n} + P_{n-2}}{Q_{n-1}a_{n-1} + \frac{Q_{n-1}}{x_n} + Q_{n-2}} = \frac{(P_{n-1}a_{n-1} + P_{n-2})x_n + P_{n-1}}{(Q_{n-1}a_{n-1} + Q_{n-2})x_n + Q_{n-1}}$$

Poiché i numeri a_n sono interi, e per ipotesi lo sono anche i numeri P_k, Q_k per $k = 0, \dots, n-1$, l'affermazione è provata ponendo:

$$P_n := P_{n-1}a_{n-1} + P_{n-2}; \quad Q_n := Q_{n-1}a_{n-1} + Q_{n-2}$$

D'altra parte, la successione dei convergenti J_n si ottiene troncando la frazione continua all'ordine n ; la formula precedente dà subito:

$$J_n = \frac{P_{n-1}a_{n-1} + P_{n-2}}{Q_{n-1}a_{n-1} + Q_{n-2}}$$

Quindi le successioni P_n e Q_n sono le successioni A_n e B_n volute perché soddisfano le medesime relazioni di ricorrenza a tre termini con le medesime condizioni iniziali.

Esercizio 6.23

Dimostrare quanto segue:

1. Le successioni P_n e Q_n dei numeratori e denominatori parziali sono successioni crescenti di numeri interi positivi;
2. I convergenti $\frac{A_n}{B_n}$ sono frazioni ridotte ai minimi termini.

Calcoliamo ora ad esempio i primi approssimanti della (6.19). Abbiamo:

$$\begin{aligned} \frac{A_1}{B_1} &= \frac{1}{2} = 0.5; & \frac{A_2}{B_2} &= \frac{1}{2 + \frac{1}{2}} = \frac{2}{5} = 0.4; & \frac{A_3}{B_3} &= \frac{1}{2 + \frac{2}{5}} = \frac{5}{12} = 0.41666\dots; \\ \frac{A_4}{B_4} &= \frac{1}{2 + \frac{5}{12}} = \frac{12}{29} = 0.4133\dots; & \frac{A_5}{B_5} &= \frac{1}{2 + \frac{12}{29}} = \frac{29}{70} = 0.428571\dots; \\ \frac{A_6}{B_6} &= \frac{1}{2 + \frac{29}{70}} = \frac{70}{169} = 0.4142011\dots; & \frac{A_{n+1}}{B_{n+1}} &= \frac{A_n a_n + A_{n-1}}{B_n a_n + B_{n-1}} \end{aligned}$$

e analogamente per la (6.20). Vediamo emergere le seguenti indicazioni numeriche, la cui validità può essere dimostrata per *una qualsiasi* frazione continua aritmetica periodica:

1. Gli approssimanti di ordine dispari $\frac{A_{2p+1}}{B_{2p+1}}$ formano una successione decrescente, mentre quelli di ordine pari $\frac{A_{2p}}{B_{2p}}$ formano una successione crescente;
2. Entrambe le successioni convergono al medesimo limite.
3. Gli approssimanti di ordine pari approssimano il limite per difetto, e quelle dispari per eccesso; dunque il troncamento all'ordine n permette anche la stima immediata dell'errore ϵ_n , definito come la differenza fra il limite e la sua approssimazione all'ordine n prefissato, cioè il termine $\frac{A_n}{B_n}$ della successione degli approssimanti. Dato che

$$\frac{A_{2p}}{B_{2p}} \leq J \leq \frac{A_{2p+1}}{B_{2p+1}}, \quad \forall p = 0, 1, 2, \dots$$

si ha evidentemente

$$\epsilon_n \leq \left| \frac{A_{n+1}}{B_{n+1}} - \frac{A_n}{B_n} \right|$$

4. Le due successioni di approssimanti pari e dispari formano due classi contigue di numeri razionali. Il loro elemento di separazione è il valore della frazione continua: ad esempio, il numero irrazionale $\sqrt{2} - 1$ per la (6.19) e il numero irrazionale $\frac{\sqrt{5} + 1}{2}$ per la (6.20).

Un teorema di Lagrange,¹⁵ che enunciamo omettendo la dimostrazione della sua parte più difficile, afferma che la corrispondenza fra numeri irrazionali quadratici e frazioni continue aritmetiche periodiche è biunivoca.

Un numero irrazionale si dice *quadratico* se è soluzione di un'equazione di secondo grado a coefficienti interi. Un numero irrazionale quadratico avrà quindi la forma $x = \frac{p}{q} + \frac{r}{s}\sqrt{n}$ dove $n \in \mathbb{N}$, $p, q, r, s \in \mathbb{Z}$, $q, s \neq 0$.

I numeri irrazionali quadratici sono casi particolari dei numeri irrazionali detti *algebrici*, cioè i numeri irrazionali che sono soluzioni di equazioni algebriche a coefficienti interi. I numeri irrazionali che non lo sono si dicono *trascendenti*.

¹⁵Giuseppe Luigi (o Joseph-Louis) Lagrange, (Torino 1736-Parigi 1813). Professore a Torino, si trasferì poi a Berlino ed infine a Parigi come Accademico di Francia e in seguito anche Professore all'École Polytechnique. Tra l'altro fu il Presidente della commissione che elaborò l'odierno sistema metrico decimale, adottato per la prima volta con la legge del 18 Germinale dell'anno III (7 aprile 1795) dalla Convenzione Nazionale della Repubblica Francese.

Ad esempio, $\sqrt{2}$ è irrazionale algebrico (Pitagora, VI secolo A.C.) perché risolve l'equazione $x^2 = 2$ mentre π è trascendente (F. von Lindemann¹⁶, 1884).

Teorema 6.5 *Un numero irrazionale ammette sviluppo in frazione continua aritmetica periodica se e solo se è quadratico.*

Dimostrazione. Limitiamoci a dimostrare solo una parte della necessità, e cioè che se x è irrazionale quadratico esso ammette sviluppo in frazione continua periodica.

Senza ridurre la generalità poniamo $x = \sqrt{n}$, $n \in \mathbb{N}$. Allora:

$$\frac{1}{\sqrt{n} + 1} = \frac{\sqrt{n} - 1}{n - 1} \Rightarrow \sqrt{n} = 1 + \frac{n - 1}{\sqrt{n} + 1}$$

Iteriamo questa espressione:

$$\begin{aligned} \sqrt{n} &= 1 + \frac{n - 1}{1 + 1 + \frac{n - 1}{\sqrt{n} + 1}} = 1 + \frac{n - 1}{2 + \frac{n - 1}{\sqrt{n} + 1}} = \dots = \\ 1 + \frac{n - 1}{2 + \frac{n - 1}{2 + \frac{n - 1}{2 + \frac{n - 1}{\sqrt{n} + 1}}}} &= \dots = 1 + \frac{n - 1}{2 + \frac{n - 1}{2 + \frac{n - 1}{2 + \frac{n - 1}{2 + \dots}}}} \end{aligned}$$

e quindi

$$\sqrt{n} - 1 = \frac{n - 1}{2 + \frac{n - 1}{2 + \frac{n - 1}{2 + \frac{n - 1}{2 + \dots}}}}$$

Il secondo membro è una frazione continua periodica di periodo 1 perché $a_k = 2$ e $b_k = n - 1$ per ogni k . Omettiamo la dimostrazione della convergenza della frazione continua, e soprattutto della sufficienza del teorema e dell'aritmeticità, cioè che il limite di ogni frazione continua periodica esiste, è un numero irrazionale quadratico, e che la frazione continua periodica può essere scelta in modo aritmetico.

¹⁶Ferdinand von Lindemann (Hannover 1852-Monaco 1939), Professore all'Università di Monaco di Baviera

Capitolo 3

Dinamica discreta

La matematica delle successioni numeriche ci consente di studiare da vicino l'evoluzione a tempo discreto (cioè ad intervalli di tempo fissati) di molti fenomeni naturali.

1 Generalità

Consideriamo l'evoluzione nel tempo di una qualsiasi quantità misurabile che siamo in grado di osservare, o che ci interessi osservare, solo ad intervalli di tempo fissati, di solito di uguale durata. Ad esempio, si può trattare della temperatura giornaliera in un determinato luogo ed ad una determinata ora, del prezzo di un titolo di borsa ogni ora, della proliferazione dei batteri in una cultura ogni minuto, della distanza dalla base di lancio di un missile ogni secondo, della portata di un fiume in un determinato luogo, cioè del volume d'acqua che vi passa ogni secondo, del numero dei giri di un motore ogni minuto, della posizione del punto luminoso sull'oscilloscopio ogni centesimo di secondo, ecc.

In ciascuno di questi casi i risultati delle osservazioni formeranno delle successioni numeriche. Potremo sempre interpretare l'indice n della successione come il tempo discreto, perché esso sarà sempre un multiplo di un intervallo fissato Δt che abbiamo la libertà di scegliere come unità di tempo: $\Delta t = 1$. Possiamo inoltre assumere che n prenda ogni valore in $\mathbb{Z}$: $n = 0$ corrisponderà all'istante iniziale, $n > 0$ al futuro e $n < 0$ al passato.

Per *dinamica discreta* si intende il problema di trovare e descrivere il comportamento della successione a_n che rappresenta l'evoluzione discreta in questione una volta assegnata la legge (fisica, biologica, economica, ecc.) che genera l'evoluzione. Per far ciò, occorre conoscere o ricavare la legge medesima e conoscere i dati iniziali. Prima di formulare il problema in

generale diamo un esempio istruttivo risolvendo un esercizio, anch'esso contenuto nel *Liber abaci* di Leonardo Fibonacci.

Esercizio 1.24

Un levriero la cui velocità cresce aritmeticamente insegue una lepre la cui velocità cresce anch'essa aritmeticamente. Quanta strada faranno prima che il levriero raggiunga la lepre?

Soluzione. Questo è un caso in cui la legge di evoluzione non è data in forma esplicita, e quindi la legge della successione a_n è incognita. Occorre anzitutto trovarla a partire dai dati del problema, e poi risolverla. Sia α_n la velocità del levriero all'istante n . Il primo dato è che essa cresce aritmeticamente. Dire che la velocità cresce aritmeticamente significa, come sappiamo, che ad ogni passo si deve aggiungere una quantità costante; quindi deve essere $\alpha_n = \alpha_{n-1} + A$ ossia $\alpha_n - \alpha_{n-1} = A$ per un qualche $A > 0$. Allo stesso modo, se β_n è la velocità della lepre all'istante n , $\beta_n - \beta_{n-1} = B$. Dobbiamo poi rendere esplicita l'ipotesi implicita nel problema che sia $0 < B < A$. Questa condizione è necessaria se si vuole che il levriero raggiunga prima o poi la lepre. La "strada" è la distanza percorsa prima del raggiungimento; ammettendo come è ovvio che la rincorsa abbia luogo su una traiettoria rettilinea, la distanza sarà un segmento (lo "spazio"). Ora la velocità è per definizione lo spazio percorso nell'unità di tempo; quindi detto s_n lo spazio all'istante $n\Delta t$ percorso dal levriero, e σ_n quello percorso dalla lepre, avremo

$$\alpha_n = \frac{s_{n+1} - s_n}{\Delta t} \quad \text{per il levriero;} \quad \beta_n = \frac{\sigma_{n+1} - \sigma_n}{\Delta t} \quad \text{per la lepre.}$$

Sostituendo $s_{n+1} - s_n$ per α_n , $s_n - s_{n-1}$ per α_{n-1} , facendo la stessa cosa con β_n e σ_n , e tenendo conto che Δt è l'unità di tempo troviamo:

$$s_{n+1} - 2s_n + s_{n-1} = A; \quad \sigma_{n+1} - 2\sigma_n + \sigma_{n-1} = B$$

Si tratta ancora di relazioni di ricorrenza a tre termini, le cui incognite sono appunto le distanze. Esse sono però diverse dal caso esaminato in (4.16) per due motivi: anzitutto non sono omogenee, e quindi non ammettono la soluzione identicamente nulla, e poi i segni del secondo e del terzo termine sono discordi. Per risolverle, osserviamo che da $\alpha_n - \alpha_{n-1} = A$ e $\beta_n - \beta_{n-1} = B$ segue subito $\alpha_n = nA + \alpha_0$ e $\beta_n = nB + \beta_0$; inoltre, inserendo in $\alpha_n = s_{n+1} - s_n$, otteniamo $s_{n+1} - s_n = nA + \alpha_0$. Quindi $s_{n+1} = s_n + nA + \alpha_0$ da cui:

$$s_1 = s_0 + \alpha_0, \quad s_2 = s_1 + A + \alpha_0 = A + 2\alpha_0 + s_0, \quad s_3 = s_2 + 2A + \alpha_0 = 3A + 3\alpha_0 + s_0, \\ s_4 = s_3 + 3A + \alpha_0 = 6A + 4\alpha_0 + s_0, \dots$$

Pertanto, in generale:

$$s_n = \frac{n(n-1)}{2}A + n\alpha_0 + s_0; \quad \sigma_n = \frac{n(n-1)}{2}B + n\beta_0 + \sigma_0, \quad \forall n \geq 2$$

Esercizio 1.25

Dimostrare per induzione le formule precedenti.

Il primo termine esprime il fatto che, se la velocità che è lo spazio percorso nell'unità di tempo cresce aritmeticamente, lo spazio percorso in n unità di tempo dovrà crescere come la somma della progressione. Sappiamo infatti che $\frac{An(n-1)}{2}$ è la somma dei primi $n-1$ termini della progressione aritmetica di ragione A .

Possiamo ora assumere, senza ledere la generalità, che tanto la lepre quanto il levriero partano da fermi, e quindi $\alpha_0 = \beta_0 = 0$. D'altra parte, la lepre parte in vantaggio: quindi $\sigma_0 > s_0$. Sia $d := \sigma_0 - s_0$ la distanza iniziale che il levriero deve colmare. Il raggiungimento avrà luogo all'intero n_0 (che corrisponde all'istante n_0 perché ricordiamo, $\Delta t = 1$) per cui risulta

$$\frac{n_0(n_0-1)}{2}A - \frac{n_0(n_0-1)}{2}B = \frac{n_0(n_0-1)}{2}(A-B) = d$$

da cui

$$n_0(A, B; d) = \frac{1}{2} \left[1 + \sqrt{1 + \frac{8d}{A-B}} \right]$$

(scartiamo ovviamente la radice negativa). n_0 non è a priori intero; potremmo prenderne la parte intera $[n_0]$, tuttavia preferiamo assumere che A, B, d siano tali che $1 + \frac{8d}{A-B} = m^2$ con m intero dispari, $m = 2p + 1$ (ipotesi questa concettualmente non restrittiva). Si ha allora $n_0 = p + 1$. Quindi la lepre percorrerà la distanza $\frac{p(p+1)}{2}B$, e il levriero la distanza $\frac{p(p+1)}{2}A + d$, dove A, B, d sono legati dalle relazioni precedenti. Si noti che:

$$\lim_{A \rightarrow B} n_0(A, B; d) = +\infty; \quad \lim_{d \rightarrow \infty} n_0(A, B; d) = +\infty; \quad \lim_{d \rightarrow 0} n_0(A, B; d) = 0$$

in accordo con tre intuizioni ovvie: la prima è il levriero non raggiungerà mai la lepre se non corre più forte di lei, la seconda è che non la raggiungerà mai se la lepre parte con un vantaggio iniziale infinito, e la terza è che la raggiungerà subito se il vantaggio iniziale della lepre è nullo. Ciò risolve l'esercizio.

Osservazione 1.21

1. Se la velocità v_n è lo spazio percorso nell'unità di tempo, $v_n = s_{n+1} - s_n$, l'accelerazione a_n , è per definizione l'incremento della velocità nell'unità di tempo: $a_n = \frac{v_{n+1} - v_n}{\Delta t} = v_{n+1} - v_n$ perché al solito prendiamo Δt come unità di tempo. Pertanto $a_n = s_{n+2} - s_{n+1} - [s_{n+1} - s_n] = s_{n+2} - 2s_n + s_n$. Quindi le ricorrenze precedenti esprimono l'ipotesi che lepre e levriero si muovono con accelerazione costante, quella del levriero essendo la più elevata delle due.

2. (*Per chi conosce già il calcolo infinitesimale*)

Sia $s(t)$ lo spazio percorso al tempo t , dove ora t è una variabile continua (cioè prende valori in $\mathbb{R}$). Allora la velocità $v(t)$ all'istante t , definita sempre come lo spazio percorso nell'unità di tempo, vale

$$\lim_{h \rightarrow 0} \frac{s(t+h) - s(t)}{h} = \frac{ds(t)}{dt}$$

(in questo contesto si usa spesso la notazione $\frac{ds(t)}{dt} = \dot{s}(t)$) e l'accelerazione, che è ancora l'incremento della velocità nell'unità di tempo, avrà l'espressione:

$$a(t) = \lim_{h \rightarrow 0} \frac{v(t+h) - v(t)}{h} = \frac{dv(t)}{dt} = \frac{d^2s(t)}{dt^2}$$

(anche qui si usa la notazione $\frac{d^2s(t)}{dt^2} = \ddot{s}(t)$). Se il tempo scorre in maniera discreta, allora possiamo porre $h = 1$, $t = nh = n$ e la formula della velocità, cioè la derivata prima, si riduce al rapporto incrementale unitario $s(n) - s(n-1)$; la indichiamo con $s_{n+1} - s_n$ e la chiamiamo *differenza prima*. Allora stesso modo, la formula dell'accelerazione, cioè la derivata seconda, si riduce a $s_{n+2} - 2s_n + s_{n-1}$ che si chiama *differenza seconda*. Infatti $v_{n+1} - v_n = s_{n+2} - s_{n+1} - [s_{n+1} - s_n] = s_{n+2} - 2s_n + s_{n-1}$. In altre parole se la variabile indipendente di una funzione (derivabile tante volte quanto occorre) viene ristretta ad assumere solo valori discreti che sono multipli interi di una quantità fissata, come avviene sempre nel calcolo numerico, le derivate diventano delle *differenze finite*.

Questo esempio mostra che la dinamica discreta consiste in una legge che si concretizza in una relazione, da ricavarsi a seconda del problema o da considerare nota, fra i primi elementi

della successione e quelli successivi. Essa descrive l'evoluzione temporale che si cerca perché l'intera successione viene ricavata iterando indefinitamente la legge in questione. Trovare la relazione temporale significa risolvere la relazione, come negli esempi del Capitolo precedente.

Definizione 1.13 *Si dice che è definita una dinamica discreta di ordine p e legge F se, data la funzione $F(x_1, \dots, x_p) : \mathbb{R}^p \rightarrow \mathbb{R}$ si pone*

$$a_{n+p} = F(a_{n+p-1}, \dots, a_n), \quad n = 0, 1, \dots \quad (1.1)$$

Qui i primi p termini della successione, detti condizioni iniziali, sono numeri reali noti:

$$a_0 = \alpha_0, a_1 = \alpha_1, \dots, a_{p-1} = \alpha_{p-1}.$$

Osservazione 1.22

1. Si assume che la F sia una funzione *regolare* dei suoi argomenti, cioè derivabile rispetto ad ogni coordinata x_k quante volte si vuole. Si ricordi poi che una funzione $f(x)$ della variabile reale x derivabile è monotona (crescente o decrescente) su un intervallo $I \subset \mathbb{R}$ se la sua derivata è ivi nonnegativa (funzione crescente) o non positiva (funzione decrescente).
2. Una dinamica di ordine p è generata da una relazione di ricorrenza a $p+1$ termini: ad esempio, la ricorrenza di Fibonacci è una dinamica discreta di ordine 2, e la relazione di ricorrenza a due termini dell'incremento o decremento esponenziale è una dinamica discreta di ordine 1.
3. La (1.1) genera una successione *diversa* per ogni diversa scelta dei dati iniziali. Assegnati i dati iniziali $\alpha_0, \dots, \alpha_{p-1}$, e usando l'abbreviazione $\alpha := (\alpha_0, \dots, \alpha_{p-1})$ l'intera successione $a_n(\alpha)$ è determinata iterando la legge F :

$$\begin{aligned} a_p(\alpha) &= F(a_{p-1}, \dots, a_0); \quad a_{p+1}(\alpha) = F(F(\alpha_{p-1}, \dots, \alpha_0), a_{p-1}, \dots, a_1) = \\ &= F(F(\alpha_{p-1}, \dots, \alpha_0), a_{p-1}, \dots, \alpha_1) \end{aligned}$$

$$a_{p+2} = F(a_{p+1}, \dots, a_2) = F(F(\alpha_p, \dots, \alpha_1), F(\alpha_{p-1}, \dots, \alpha_0), \dots, \alpha_2)$$

Definendo la legge di composizione $F \circ F$ nel seguente modo naturale:

$$\begin{aligned} F(F(\alpha_{p-1}, \dots, \alpha_0), \dots, \alpha_1) &= (F \circ F)(\alpha_{p-1}, \dots, \alpha_0); \\ F(F(\alpha_p, \dots, \alpha_1), F(\alpha_{p-1}, \dots, \alpha_0), \dots, \alpha_2) &= (F \circ F \circ F)(\alpha_{p-1}, \dots, \alpha_0) \end{aligned}$$

otteniamo

$$\alpha_{p+n} = (F \circ F \dots \circ F)(\alpha_{p-1}, \dots, \alpha_0) := F^n(\alpha_{p-1}, \dots, \alpha_0) \quad (1.2)$$

dove l'ultima abbreviazione sta per la n -sima composizione della funzione F e non per una elevazione a potenza.

4. Il caso più semplice si ha per $p = 1$. Qui, data una funzione regolare $f(x) : I \subset \mathbb{R} \rightarrow \mathbb{R}$, dove I è un intervallo chiuso qualsiasi, e il codominio deve contenere l'intervallo I medesimo, la dinamica è semplicemente definita dalla formula della composizione abituale successiva

$$\begin{aligned} x_{n+1} &= f(x_n) = f^n(x_0); & n = 0, 1, \dots \\ f^n(x_0) &= f(x_n) = f(f(x_{n-1})) = f(f(f(x_{n-2}))) = \dots = f(\dots f(x_0) \dots) \end{aligned} \quad (1.3)$$

dove al solito il dato iniziale è da considerare noto.

Ricordiamo che un intervallo $I \subset \mathbb{R}$ di estremi $a < b$ *limitato* (cioè $-\infty < a < b < +\infty$) si dice *chiuso* se contiene i suoi estremi: $I = \{x \in \mathbb{R} \mid a \leq x \leq b\}$. Si denota di solito $[a, b]$ un intervallo chiuso di estremi $a < b$. Un intervallo si dice *aperto* se non contiene i suoi estremi, e lo si denota $]a, b[$. In altre parole $]a, b[:= \{x \in \mathbb{R} \mid a < x < b\}$. Analogamente, $]a, b]$ denota un intervallo *aperto a sinistra*, e $[a, b[$ un intervallo *aperto a destra*.

Esercizio 1.26

Siano I_1 e I_2 intervalli limitati, con $I_1 \subset I_2$. Sia $J := I_2 \setminus I_1$ l'insieme complementare di I_1 in I_2 .

Se f è una funzione lineare, $f = \beta x$, la legge (1.3) si riconduce alla relazione di ricorrenza a due termini studiata nel Cap.1 che dà l'incremento o il decremento esponenziale: $x_{n+1} = \beta^n x_0$.

Altri esempi concreti se ne possono fare quanti se ne vogliono:

1. $f(x) = \sin x$ definita su $I = [0, 2\pi]$. Sia $x_0 = \pi/2$. Allora $x_1 = 1$, $x_2 = \sin 1$, $x_3 = \sin(\sin 1)$, $x_4 = \sin(\sin(\sin 1))$, $\dots$

2. $f(x) = e^x$ definita su $\mathbb{R}$. Sia $x_0 = 0$. Allora $x_1 = 1$, $x_2 = e^1 = e$, $x_3 = e^e$, $x_4 = e^{e^e}$;
3. $f(x) = x^2 + c$ definita su $\mathbb{R}$; c è una costante reale. Sia $x_0 = 0$. Allora $x_1 = c$, $x_2 = c^2 + c$, $x_3 = c^2(c+1)^2 + c, \dots$
5. Riconsideriamo infine la mappa logistica: $f(x) = kx(1-x)$. Si ha:

$$\max_{0 \leq x \leq 1} x(1-x) = \frac{1}{4}$$

perché il massimo di $g(x) := x(1-x)$ su $[0, 1]$ è raggiunto a $x = 1/2$ per evidenti ragioni di simmetria (per chi conosce il calcolo infinitesimale: $g'(1/2) = 0$, $g''(1/2) = -2 < 0$; $g(0) = g(1) = 0$). Allora dobbiamo richiedere $k \leq 4$ affinché il codominio della mappa sia contenuto in $[0, 1]$ in modo da poterla iterare. Sia $x_0 = \frac{1}{2}$. Allora $x_1 = k\frac{1}{4}$, $x_2 = k\frac{1}{4}(1 - k\frac{1}{4}), \dots$

6. Se la funzione f è iniettiva possiamo definire la dinamica anche per n negativo nel modo seguente:

$$x_{-1} = f^{-1}(x_0), \quad x_{-2} = f^{-1}(x_{-1}) = f^{-1}(f^{-1}(x_0)), \dots, x_{-n-1} = f^{-n}(x_0) \quad (1.4)$$

dove al solito $f^{-1}(x)$ denota la funzione inversa di $f(x)$: $(f^{-1} \circ f)(x) = (f \circ f^{-1})(x) = x$. In questo caso potremo distinguere, come abbiamo già fatto più volte, fra l'evoluzione *in avanti*, o *nel futuro*, che avviene per definizione se $n > 0$, e l'evoluzione *all'indietro*, o *nel passato*, che avviene per definizione se $n < 0$.

Sia ad esempio $f(x) = e^x$ definita su $\mathbb{R}$. Poiché $f : \mathbb{R} \rightarrow \mathbb{R}_+$ è iniettiva $f^{-1}(x) = \ln x$ è definita per $x \in \mathbb{R}_+$. Allora $x_{-1} = \ln x_0$, $x_{-2} = \ln \ln x_0$, $\dots$, $x_{-n-1} = \ln \ln \dots \ln x_0$. Con la convenzione naturale $f^0(x) = I_d$ si vede subito che $f^n \circ f^{-n}(x) = f^{n-n} \circ f^n(x) = x \quad \forall x \in \mathbb{R}_+$.

Esempio 1.17

1. Le relazioni di ricorrenza a due, tre, p termini studiate nel Capitolo 1 precedente sono ovviamente tutti esempi di dinamica discreta. Esse hanno la particolarità, già rilevata, di essere *lineari*: combinazioni lineari di soluzioni sono ancora soluzioni. Questa particolarità è dovuta al fatto che la funzione F della (1.1) è in questi casi essa stessa

lineare, cioè un polinomio di primo grado nelle $a_{n+p-1}, \dots, a_n$. Ad esempio per il caso di Fibonacci $a_{n+2} = a_{n+1} + a_n$ si ha $F(a_{n+1}, a_n) = a_{n+1} + a_n$ che è una funzione di primo grado. Abbiamo poi visto come esse ammettano soluzione unica una volta assegnati i dati iniziali, e le abbiamo determinate esplicitamente.

2. Consideriamo ora una "legge" quadratica, e precisamente la relazione di ricorrenza a due termini di tipo quadratico

$$x_{n+1} = F(x_n); \quad F(x) := \mu x(1 - x), \quad x \in [-1, 1], \quad 0 < \mu < 4 \quad (1.5)$$

Essa è nota col nome di "mappa logistica". C'è un solo dato iniziale, $x_0 = p \in [-1, 1]$. Si noti che possiamo riscrivere la legge (1.5) anche sotto la forma $x_{n+1} = F \circ F \dots \circ F(x_0) := (\circ F)^n(x_0)$ dove al solito il simbolo $f \circ g$ significa la composizione della funzione f con la funzione g .

3. Si parla di dinamica discreta perché questa nozione altro non è che la formulazione in astratto delle nozioni più semplici della dinamica assumendo di osservare il flusso del tempo ad intervalli fissati. Ad esempio, esaminiamo il significato meccanico dell'esempio della lepre e del levriero. L'accelerazione è per definizione l'incremento della velocità nell'unità di tempo. Quindi se la velocità cresce in progressione aritmetica la differenza fra le velocità a due istanti successivi è costante. Ciò equivale ad affermare che l'accelerazione è costante. I moti ad accelerazione costante si dicono uniformemente accelerati, e tali dunque sono quelli della lepre e del levriero. La maggiore accelerazione del levriero gli consente di aumentare progressivamente la velocità più della lepre e quindi di raggiungerla.
4. (*Per chi conosce già il calcolo infinitesimale*). La legge elementare del moto uniformemente accelerato afferma che la distanza percorsa dal mobile trascorso un tempo t è proporzionale a t^2 . Possiamo ricavarla tramite il calcolo infinitesimale. Se $a = \frac{dv}{dt} = A$, allora $v(t) = At + B$. Dato poi che $v = \frac{ds}{dt}$, da $\frac{ds}{dt} = At + B$ ricaviamo $s(t) = \frac{1}{2}t^2 + Bt + C$; qui la ritroviamo al trascorrere del tempo discreto perché dopo n "secondi" la distanza percorsa è proporzionale a $n(n+1)$. La seconda legge della dinamica afferma che l'accelerazione è proporzionale alla forza impressa al mobile: possiamo così dire che

la legge del moto di lepre e levriero è la stessa che avrebbero se fossero sottoposti all'azione di una forza costante (ad esempio, la gravità).

2 Orbite, punti fissi, periodicità

Consideriamo la legge di evoluzione (1.1). La domanda fondamentale alla quale la dinamica discreta (come del resto qualsiasi dinamica) deve rispondere è la seguente:

Scelto arbitrariamente un dato iniziale, quale sarà la sua evoluzione?

In altre parole, quale sarà la successione generata dalla (1.1) in corrispondenza ad ogni dato iniziale? Per una scelta generale di F sarà impossibile dare una risposta completa a questa domanda risolvendo esplicitamente la ricorrenza come abbiamo fatto per gli esempi in precedenza. Anche per il caso più semplice di mappa quadratica, la mappa logistica (1.5), ciò risulta impossibile. Per la soluzione quantitativa ci si può affidare al calcolatore; tuttavia bisogna imparare a porre a questo strumento le domande giuste per non usarlo in modo bovino. (L'uso del calcolatore in modo bovino, cioè acritico, spesso ci confonde le idee e addirittura ci conduce in grave errore invece di insegnarci qualcosa. In particolare la dipendenza dai dati iniziali può essere così delicata da rendere praticamente impossibile il calcolo delle traiettorie, come discuteremo in seguito). In altre parole si deve fare prima uno studio qualitativo dell'evoluzione. Questa locuzione significa che si vogliono capire a priori le proprietà più rilevanti dell'evoluzione senza risolvere esplicitamente la ricorrenza. A questo scopo occorre elaborare gli opportuni concetti matematici. Fatto ciò, avremo le idee sufficientemente chiare per sapere cosa dovremo effettivamente esaminare con l'aiuto del calcolatore.

Cominciamo quindi col dare alcune definizioni fondamentali relativi alle successioni generate dalla (1.1). Osserviamo anzitutto che poiché i dati iniziali $\alpha_0, \dots, \alpha_{p-1}$ sono numeri reali qualsiasi possiamo senz'altro identificare l'insieme di tutti i dati iniziali con lo spazio vettoriale $\mathbb{R}^p$ a p dimensioni, $p \geq 1$. Nel seguito useremo la notazione vettoriale, cioè scriveremo $\alpha := (\alpha_0, \dots, \alpha_{p-1})$, in corrispondenza alla decomposizione cartesiana lungo la base canonica $\alpha = \alpha_0 e_1 + \dots + \alpha_{p-1} e_p$. Qui al solito $e_1 = (1, 0, \dots, 0)$; $e_2 = (0, 1, \dots, 0)$; ..., $e_p = (0, 0, \dots, 1)$. Ad esempio, per $p = 0$ (la retta) avremo semplicemente il dato iniziale $\alpha \in \mathbb{R}$, cioè un punto sulla retta; per $p = 1$ avremo i due dati iniziali $(\alpha_0, \alpha_1) \in \mathbb{R}^2$, identificabili con un punto nel piano cartesiano, per $p = 2$ tre dati iniziali $(\alpha_0, \alpha_1, \alpha_2) \in \mathbb{R}^3$,

identificabili con un punto nello spazio ecc.

Definizione 2.14 Sia $a_n(\alpha)$ la successione generata dalla (3.8) in corrispondenza ad un fissato dato iniziale $\alpha \in \mathbb{R}^p$. Allora

1. L'insieme numerico $a_n(\alpha) : n \in \mathbb{Z}$ dei valori assunti dalla successione si dice *orbita del dato iniziale* α .
2. Se esiste $N \in \mathbb{N}, N \geq 2$ tale che $a_n(\alpha) = a_{n+N}(\alpha), \forall n = 0, 1, 2, \dots$ l'orbita consiste esattamente di N punti e si dice *periodica di periodo* N .
3. Un'orbita che consiste di un punto solo (e quindi periodica di periodo 1) si dice *punto fisso, o punto di equilibrio*.

Esempio 2.18

1. Prendendo l'unione su $n \in \mathbb{N}$ dei valori dei numeri di Fibonacci (4.15) si ottiene l'orbita di dato iniziale $(0, 1)$ generata dalla legge di ricorrenza a tre termini $a_{n+2} = a_{n+1} + a_n$; prendendo quella dei valori della soluzione (4.17) si ottiene l'orbita di ogni dato iniziale $(p, q) \in \mathbb{R}^2$ generata dalla ricorrenza a tre termini $A_{n+2} = aA_{n+1} + bA_n$. Ciascuna di queste "leggi" è ovviamente un caso particolare della (3.8): basta prendere $F = x_1 + x_2$ nel primo caso e $F = ax_1 + bx_2$ nel secondo. In ognuno di questi casi i punti dell'orbita sono tutti distinti.
2. Consideriamo invece ancora la ricorrenza $a_{n+2} + a_n = 0$ dell'osservazione 4.5 del Capitolo 1. Per ogni dato iniziale $(p, q) \in \mathbb{R}^2$ si ha $a_{2n} = (-1)^n p, a_{2n+1} = (-1)^n q$ per cui l'orbita è composta esattamente dai quattro punti $\pm p$ e $\pm q$. Se $p = q = 1$ addirittura dai soli due punti ± 1 . Quindi è sufficiente che un'orbita contenga anche solo due punti distinti per non coincidere con un punto fisso. Nel primo caso ogni orbita è periodica di periodo 4; nel secondo caso si tratta di un'orbita di periodo 2.
3. L'origine in $\mathbb{R}^p$, cioè il punto di coordinate $(0, 0, \dots, 0)$, è ovviamente un punto fisso per l'evoluzione generata dalla ricorrenza lineare omogenea a p termini

$$a_{n+p} + b_1 a_{n+p-1} + \dots + b_n a_n = 0 \quad (2.6)$$

- $(b_1, \dots, b_n)$ sono delle costanti). L'orbita generata da queste condizioni iniziali è infatti la successione nulla $(0, 0, \dots)$ che prende solo il valore 0. In particolare, l'origine è pertanto un punto fisso per tutte le successioni del tipo di Fibonacci studiate finora.
4. Consideriamo nella (2.6) il dato iniziale $\alpha_0 = \dots \alpha_{p-1} = \alpha$. Allora, se la (2.6) medesima ammette la soluzione $\alpha_p = \alpha$, α è chiaramente un punto fisso. Esempio: la legge del moto uniformemente accelerato dell'esercizio precedente con accelerazione nulla, cioè $x_{n+2} - 2x_{n+1} + x_n = 0$. Sia $x_0 = x_1 = \alpha$, $\forall \alpha \in \mathbb{R}$. Allora $x_2 = \alpha$ risolve la ricorrenza e quindi $\forall \alpha \in \mathbb{R}$ è un punto fisso.
5. Consideriamo la mappa quadratica (1.5). Allora il punto $x = 0$ è chiaramente punto fisso perché $F(x) = \mu x(1 - x) = 0$ per $x = 0$ e quindi $x_n = 0 \forall n = 0, 1, \dots$. Il punto $x = 1$ non è punto fisso anche se $F(1) = 0$. L'orbita è infatti $(1, 0, 0, \dots)$ ed è composta da due punti. Punti iniziali come questo, che evolvono in punti fissi verranno in seguito definiti come punti successivamente fissi.

Osservazione 2.23

Sia $\mathcal{G}$ un gruppo ciclico generato dall'azione dell'operazione G . In algebra si definisce *orbita* dell'elemento $g \in \mathcal{G}$ l'insieme di tutti gli elementi $G^n g : n \in \mathbb{Z}$. Questa nozione è la versione astratta del concetto di orbita introdotto in precedenza purché la dinamica sia autonoma. Per verificare questa affermazione dobbiamo fare vedere che le orbite qui definite formano gruppo per ogni scelta del dato iniziale. Consideriamo per semplicità solo il caso unidimensionale $p = 1$. La dinamica discreta è generata per composizioni successive dalla funzione $f(x)$, che supponiamo iniettiva. Dato $x_0 \in I$, mostriamo che l'insieme $U := \bigcup_{n \in \mathbb{Z}} f^n(x_0)$ forma un gruppo abeliano rispetto all'operazione $\circ f$, la composizione con la funzione f , nella convenzione naturale che $f^{(0)}$ sia l'identità. Si ha chiaramente

$$f^{m+n}(x_0) = f^n \circ (f^m(x_0)) = f^m \circ (f^n(x_0)) \quad \forall (m, n) \in \mathbb{Z}$$

e per $m = -n$ otteniamo $f^{-n} \circ (f^n(x_0)) = f^n \circ (f^{-n}(x_0)) = x_0$. Si tratta chiaramente del gruppo ciclico generata da f rispetto all'operazione di composizione di funzioni. La dinamica è autonoma nel senso che la forma della funzione f è costante rispetto all'iterazione (non dipende cioè da n).

3 Caratterizzazione dei punti fissi

Per studiare l'andamento delle orbite generate dalla dinamica discreta cominciamo dal caso più semplice, quello delle orbite composte da un solo punto: i punti fissi, o punti di equilibrio. Questa locuzione viene dall'analogia meccanica: un punto che non si muove sotto l'azione di una dinamica (necessariamente generata da una forza) si dice in equilibrio, o in quiete. Il principio d'inerzia di Galileo afferma che un corpo non soggetto a forze o sta in quiete o si muove di moto rettilineo uniforme, cioè con velocità costante. Anche qui cominciamo dal caso più semplice, $p = 1$, la dinamica discreta generata dall'iterazione di una mappa unidimensionale $f(x) : I \rightarrow \mathbb{R}$. L'affermazione seguente è una conseguenza diretta della definizione di punto fisso:

Proposizione 3.2 *Sia $I \subset \mathbb{R}$ un intervallo non necessariamente limitato. I punti fissi della dinamica discreta generata dall'iterazione della funzione $f(x) : I \rightarrow \mathbb{R}$ sono tutte e sole le soluzioni dell'equazione $x = f(x)$, $x \in I$.*

Esempio 3.19

1. Consideriamo ancora la mappa logistica $f(x) = \mu x(1 - x)$, $x \in I = [0, 1]$, $\mu > 0$. I suoi punti fissi sono le soluzioni dell'equazione $x = \mu x(1 - x)$. Il primo punto fisso è il punto $x_1 = 0$ come già sappiamo; il secondo zero dell'equazione è $x_2 = \frac{\mu - 1}{\mu}$, che è accettabile come punto fisso se $|\mu - 1| \leq \mu$. D'ora in poi supponiamo $\mu > 1$ cosicché possiamo concludere che la mappa logistica ammette i due punti fissi $x_1 = 0$ e $x_2 = \frac{\mu - 1}{\mu}$.
2. Consideriamo le mappe $f_1(x) = x^3$ e $f_2(x) = x^2$ definite su $\mathbb{R}$. f_1 ammette i tre punti fissi $x_1 = 0$ e $x_{2,3} = \pm 1$. La mappa f_2 ammette il punto fisso $x = 1$; invece il punto $x = -1$ non è fisso perché dà origine alla successione $-1, 1, \dots, 1, \dots$

Quest'ultimo esempio, assieme all'orbita generata dal punto $x = 1$ della mappa logistica, motiva la seguente generalizzazione della nozione di punto fisso:

Definizione 3.15 *Un punto $\alpha = (\alpha_1, \dots, \alpha_p)$ si dice successivamente fisso rispetto alla dinamica discreta F definita dalla (3.8) se esiste $N \in \mathbb{N}$ tale che $a_n(\alpha) = a_N(\alpha) \forall n \geq N$.*

Esempio 3.20

Si denoti S^1 la circonferenza di raggio 1. Ogni punto di S^1 può essere individuato tramite una coordinata angolare θ , $0 \leq \theta \leq 2\pi$. L'angolo non cambia se gli aggiungiamo un multiplo intero di 2π : $\theta \equiv \theta + 2k\pi$, $k \in \mathbb{Z}$. Quindi possiamo anche scrivere $S^1 = \mathbb{R} \bmod 2\pi$. Ora la funzione $f(\theta) : \theta \mapsto 2\theta$ è ben definita da S^1 in S^1 perché $f(\theta + 2k\pi) = 2\theta + 4k\pi = 2\theta = f(\theta)$. Si noti per di più che f è lineare per cui $(f(\theta))^n = f^n(\theta)$. Tuttavia non è biunivoca. Anzi, è un'applicazione 2 a 1. Infatti tutti i punti y e $\pi + y$, $0 \leq y < \pi$, che sono distinti, vengono inviati nel medesimo punto $2y$. Troviamo ora i punti fissi ed i punti successivamente fissi per la dinamica in S^1 definita da $f(\theta)$. $\theta = 0$ è ovviamente un punto fisso. Sia ora $\theta_N := \frac{2k\pi}{2^N}$. Allora $f^N(\theta_N) = 2k\pi \equiv 0$ e quindi tutti i punti θ_N sono successivamente fissi. L'analogia meccanica è la seguente: ad ogni "secondo" il punto θ compie una rotazione di ampiezza pari a θ stesso; allora il punto $\frac{2k\pi}{2^N}$ impiega k secondi per arrivare all'origine e lì si ferma.

Si noti infine che fra tutte le applicazioni lineari $\theta \mapsto A\theta$ quelle e solo quelle per cui $A \in \mathbb{Z}$ definiscono un'applicazione di S^1 in sè: in tal caso infatti, e solo in questo, $A(\theta + 2k\pi) = A\theta + 2Ak\pi = A\theta$.

Nel caso generale della dinamica discreta di ordine p consideriamo per il momento solo il caso più semplice, quello lineare, in cui la (3.8) assume la forma

$$a_{n+p} + b_{p-1}a_{n+p-1} + \dots + b_0a_n = b_p \quad (3.7)$$

con le condizioni iniziali $a_i = \alpha_i$, $i = 0, \dots, p-1$. Se $b_p = 0$ la ricorrenza (3.7) si dice *omogenea*. Qui possiamo formulare in via generale alcune delle osservazioni del paragrafo precedente:

Proposizione 3.3

1. Nel caso omogeneo $b_p = 0$ la soluzione banale $\gamma = 0$ è un punto fisso.
2. Sia $(1 + b_{p-1} + \dots + b_0) \neq 0$. Allora $\gamma := \frac{b_p}{1 + b_{p-1} + \dots + b_0}$ è un punto fisso per la dinamica di ordine p generata dalla (3.7). In particolare: la soluzione banale $\gamma = 0$ è sempre un punto fisso nel caso omogeneo, e non lo è mai nel caso non omogeneo.
3. Sia $1 + b_{p-1} + \dots + b_0 = 0$. Allora la (3.7) non ha punti fissi nel caso non omogeneo; ogni $\gamma \in \mathbb{R}$ è fisso nel caso omogeneo.

Dimostrazione. L'affermazione 1 è evidente; per verificare la 2 basta sostituire l'espressione di γ nella (3.7) nel caso non omogeneo, e il fatto che $\gamma = 0$ non sia punto fisso nel caso non omogeneo è banale. Quanto all'affermazione 3, se $1 + b_{p-1} + \dots + b_0 = 0$ nessun punto $\gamma \in \mathbb{R}$ può essere un'orbita nel caso non omogeneo perché il primo membro di (3.7) diventa $\gamma(1 + b_{p-1} + \dots + b_0) = 0 \neq b_p$; nel caso omogeneo invece ciascun γ lo è proprio perché $\gamma(1 + b_{p-1} + \dots + b_0) = 0$. Ciò prova la proposizione.

Esempio 3.21

Riprendiamo la dinamica discreta della lepre e del levriero, cioè la ricorrenza a tre termini $a_{n+2} - 2a_{n+1} + a_n = 1$. Qui $p_1 = -2, p_2 = 1$ e quindi $1 + p_1 + p_2 = 0$. Dunque non ci sono punti fissi, e d'altra parte ogni punto è fisso nel caso omogeneo. L'interpretazione meccanica di questo esempio è istruttiva: la quantità $a_{n+2} - 2a_{n+1} + a_n$ è l'accelerazione discreta. Se essa è costante, ma non nulla, nessun punto potrà rimanere fermo, e quindi non potrà esserci alcun punto fisso perché i punti fissi sono punti di equilibrio; se invece è inizialmente nulla, insieme alla velocità $a_n - a_{n-1}$, non potrà esserci alcuna evoluzione perché un punto a velocità nulla, cioè fermo, deve possedere accelerazione non nulla per acquistare velocità e potersene andare. (In particolare, lepre e levriero staranno sempre fermi alla medesima distanza).

Nel caso generale della (3.8) vale un'affermazione analoga alla Proposizione 3.2. La si può enunciare così:

Proposizione 3.4 *I punti fissi della dinamica discreta generata dall'iterazione della $F(x_1, \dots, x_p) : \mathbb{R}^p \rightarrow \mathbb{R}$ cioè dall'equazione*

$$a_{n+p} = F(a_{n+p-1}, \dots, a_n), \quad n = 0, 1, \dots \quad (3.8)$$

sono le soluzioni delle equazioni

$$\begin{cases} a_0 = \dots = a_{p-1} := x \\ x = F(x, x, \dots, x) \end{cases}$$

4 Punti periodici e orbite periodiche

Definizione 4.16 *Un punto $\alpha = (\alpha_1, \dots, \alpha_p) \in \mathbb{R}^p$ che preso come dato iniziale per la dinamica discreta di ordine p (3.8) dà origine ad un'orbita periodica si dice punto periodico*

per la dinamica F . Il più piccolo intero N per cui $F^N(\alpha) = \alpha$ si dice periodo primitivo. Un punto periodico con periodo primitivo 1 è un punto fisso.

Osservazione 4.24

Sia $p = 1$. Allora $\alpha \in \mathbb{R}$ è periodico se esiste $N \in \mathbb{N}$ tale che $F^N(\alpha) = \alpha$.

Esempio 4.22

1. Consideriamo la ricorrenza $a_{n+2} + a_n = 0$ dell'Osservazione 4.5 del Capitolo 1. Si tratta come abbiamo visto un caso particolare della (3.8) con $p = 2$ e $F(a_{n+1}, a_n) = -a_n$. Allora ogni dato iniziale $\alpha = (\alpha_1, \alpha_2) \in \mathbb{R}^2$ è periodico. Sia infatti $a_0 = \alpha_0$; $a_1 = \alpha_1$. Allora $a_2 = -\alpha_0$, $a_3 = -\alpha_1$, $a_4 = \alpha_0$, $a_5 = \alpha_1$ ed in generale $a_{2n} = (-1)^n a_0$, $a_{2n+1} = (-1)^n a_1$.
2. Consideriamo il caso $p = 1$. La dinamica lineare $a_{n+1} = \lambda a_n$ (che corrisponde chiaramente alla funzione $F(a) = \lambda a$), $\lambda \in \mathbb{R}$, non ha ovviamente alcun punto periodico. Sia poi $F(a) = a^3$ definita su $\mathbb{R}$. È chiaro che F non ha punti periodici, perché qui $F^n(a) = a^{3^n}$, e le soluzioni reali dell'equazione $a^{3^n} = a$ sono $a = 0$, $a = \pm 1$ che sono punti fissi.
3. Sempre con $p = 1$ sia $F(a) = a^2 - 1$, $a \in \mathbb{R}$. I punti $a = 0$ e $a = 1$ sono periodici, e generano orbite di periodo 2. Infatti $F(0) = -1$, $(F \circ F)(0) = 0$, $(F \circ F \circ F)(0) = -1, \dots$. D'altra parte, $F(1) = 0$, $(F \circ F)(1) = 1$, $(F \circ F \circ F)(1) = 0, \dots$.

4. Esercizio

Dimostrare che $F^{2n}(0) = 1$, $F^{2n+1}(0) = 0$, $n = 0, 1, 2, \dots$

5. Riprendiamo la mappa $F(\theta) = 2\theta$ della circonferenza S^1 in sè. Anche qui $F^n(\theta) = 2^n \theta$. Pertanto il dato iniziale θ sarà periodico se e solo se $2^n \theta = \theta + 2k\pi$ per un qualche $n \geq 2$ ed un qualche intero k . Dunque tutti i punti θ del tipo

$$\theta = \frac{2k\pi}{2^n - 1}, \quad n = 2, \dots, \quad 0 \leq k \leq 2^n - 1$$

sono periodici.

Un esempio molto significativo di dinamica discreta è la traslazione della circonferenza.

Definizione 4.17 *Dicesi traslazione della circonferenza di angolo $\lambda \in [0, [2\pi$ la mappa $\theta \mapsto T_\lambda(\theta)$ da S^1 in sé definita nel modo seguente:*

$$T_\lambda(\theta) = \theta + 2\pi\lambda \quad (4.9)$$

Osservazione 4.25

La mappa T_λ , essendo una traslazione, è biunivoca e lascia invariate le lunghezze degli archi sulla circonferenza. Consideriamo infatti l'arco Γ di estremi $\alpha < \beta$. La sua lunghezza è $l(\Gamma) = \beta - \alpha$. Calcoliamo la lunghezza $l(T_\lambda\Gamma)$ dell'immagine di Γ attraverso T_λ . Si ha, per definizione di lunghezza d'arco:

$$l(T_\lambda\Gamma) = T_\lambda\beta - T_\lambda\alpha = \beta - 2\pi\lambda - \alpha + 2\pi\lambda = \beta - \alpha = l(\Gamma)$$

come volevasi dimostrare. Inoltre l'inversa è evidentemente $T_\lambda^{-1}(\theta) = \theta - 2\pi\lambda$.

La mappa T_λ ha un comportamento radicalmente diverso a seconda della natura aritmetica di λ . Il teorema seguente, molto famoso, è uno dei modi migliori per mettere in risalto la natura differente fra numeri razionali e irrazionali.

Teorema 4.6 *Per la mappa T_λ da S^1 in sé definita dalla (4.9) vale la seguente alternativa:*

1. *Se $\lambda \in \mathbb{Q}$, $\lambda = \frac{p}{q}$, tutti i punti $\theta \in S^1$ sono periodici di periodo q , cioè $[T_{p/q}(\theta)]^q = \theta$.*
2. *(Teorema di Jacobi¹). Se $\lambda \notin \mathbb{Q}$ tutti i punti di ogni orbita $T_\lambda^n(\theta) : n \in \mathbb{Z}, \theta \in S^1$ sono distinti e per ogni dato iniziale θ l'orbita è densa in S^1 , cioè: se $\phi \in [0, [2\pi[$ è un punto qualsiasi di S^1 , $\forall \epsilon > 0$ esiste $n(\epsilon) \in \mathbb{Z}$ tale che $|T_\lambda^{n(\epsilon)} - \phi| < \epsilon$.*

Dimostrazione. Sia $\lambda < \frac{p}{q}$. Allora, applicando la definizione (4.9):

$$(T_{p/q}(\theta))^q = \theta + 2\pi q \frac{p}{q} = \theta + 2p\pi = \theta$$

e ciò prova la prima affermazione. Quanto alla seconda, ci limitiamo a dimostrare che per ogni θ i punti dell'orbita sono tutti distinti. Supponiamo per assurdo $T_\lambda^n(\theta) = T_\lambda^m(\theta)$ con

¹Carl Gustav Jacobi (Potsdam 1804-Berlino 1851), Professore di Matematica alle Università di Königsberg e poi di Berlino.

$n \neq m$. Dato che $T_\lambda^n(\theta) - T_\lambda^m(\theta) = 2\pi(n-m)\lambda$ si concluderebbe che $2\pi(n-m)\lambda = 0 \bmod 2\pi$ cioè $(n-m)\lambda \in \mathbb{Z}$ il che contraddice l'ipotesi. In particolare, se λ è irrazionale nessuna condizione iniziale può dare origine ad un punto fisso o ad un punto periodico. L'orbita è quindi composta da infiniti punti.

Osservazione 4.26

1. Si può dimostrare (teorema di Weyl²) che questi punti non solo sono densi sulla circonferenza, ma che formano anche una successione *uniformemente ripartita*. Ciò significa quanto segue: dato l'arco di circonferenza S_α di apertura α , e dato $N > 0$ sia f_α^N la *frequenza di visita* del settore S_α della successione dei primi N termini di $T_\lambda^n(\theta)$:

$$f_\alpha^N := \frac{\{\#n : T_\lambda^n(\theta) \in S_\alpha \mid 1 \leq n \leq N\}}{N}$$

In parole: f_α^N è la frazione dei primi N termini della successione che va a cadere nell'arco S_α . Allora la successione è uniformemente ripartita se $\lim_{N \rightarrow \infty} f_\alpha^N = \alpha$. In parole: al tendere di n all'infinito i punti della successione cadono in ogni dato arco di circonferenza in proporzione alla sua lunghezza.

2. Esaminiamo il significato meccanico della mappa T_λ : all'angolo θ essa associa l'angolo $\theta + 2\pi\lambda$. Equivalentemente, al punto di coordinate $x = \cos(\theta), y = \sin(\theta)$ sulla circonferenza unitaria viene associato il punto di coordinate $x = \cos(\theta + 2\pi\lambda), y = \sin(\theta + 2\pi\lambda)$. Iterando, avremo che la mappa T_λ^n associa al punto $x = \cos(\theta), y = \sin(\theta)$ su S^1 il punto di coordinate $x_n = \cos(\theta + 2n\pi\lambda), y_n = \sin(\theta + 2n\pi\lambda)$. Se al solito interpretiamo la successione n come il trascorrere del tempo di un quantità discreta, ad esempio di un secondo, il punto (x_n, y_n) sarà la posizione raggiunta dopo n secondi dal dato iniziale (x, y) a seguito del moto circolare uniforme di velocità angolare $2\pi\lambda$. L'insieme dei punti $(x_n, y_n) : n \in \mathbb{Z}$ sarà l'orbita di questo moto circolare uniforme. Il teorema di Weyl afferma dunque che l'orbita del moto circolare uniforme è uniformemente ripartita sulla circonferenza per ogni dato iniziale se la velocità angolare è irrazionale.

²Hermann Weyl (Amburgo 1885-Zurigo 1955), Professore al Politecnico Federale di Zurigo, poi a Berlino, Göttingen e infine a Princeton.

5 Punti fissi. Stabilità, instabilità, attrazione, repulsione

Come abbiamo già ricordato, lo scopo della dinamica discreta è la descrizione di tutte le orbite possibili, cioè la risposta alla domanda: scelto arbitrariamente un dato iniziale, come sarà fatta l'orbita? Vorremmo cioè essere capaci di determinare se (e come) la dinamica discreta a partire da un dato iniziale scelto a piacere darà origine ad un punto fisso o ad un'orbita periodica, o comunque determinare come si comporterà l'orbita per $n \rightarrow \infty$. Poiché i soli dati iniziali per i quali sappiamo descrivere completamente l'orbita sono i punti fissi e quelli periodici, è chiaro che la prima cosa da fare sarà esaminare l'evoluzione dei dati iniziali che si discostano poco da questi, dove la locuzione "discostarsi" poco andrà precisata matematicamente.

Vediamo alcuni esempi molto semplici.

Esempio 5.23

Consideriamo la ricorrenza a due termini (equivalentemente, mappa lineare di $\mathbb{R}$ in sé) $a_{n+1} = \lambda a_n$, con la condizione iniziale $a_0 = \alpha \in \mathbb{R}$. Sappiamo già che la soluzione è $a_n = \lambda^n x$, $n \in \mathbb{Z}$ (futuro per $n > 0$, passato per $n < 0$). Qui c'è un solo punto fisso $\forall \lambda$, $x = 0$; per $\lambda = 1$ ogni punto è fisso, e per $\lambda = -1$ ogni punto è periodico di periodo 2. Sia pertanto $|\lambda| \neq 1$. Distinguiamo ora quattro casi:

1. $0 < \lambda < 1$. In tal caso, come sappiamo, $\lim_{n \rightarrow +\infty} a_n = \lim_{n \rightarrow +\infty} \lambda^n x = 0$, $\forall x \in \mathbb{R}$.
2. $\lambda > 1$. Anche questo caso è già stato esaminato, e sappiamo che $\lim_{n \rightarrow +\infty} a_n = \lim_{n \rightarrow +\infty} \lambda^n x = \infty$.
3. $-1 < \lambda < 0$. Qui abbiamo $a_{2p} = |\lambda|^{2p} x \rightarrow 0$ per $p \rightarrow \infty$ decrescendo, cioè da valori positivi, mentre $a_{2p+1} = -|\lambda|^{2p+1} x \rightarrow 0$ crescendo, cioè da valori negativi. Quindi $\forall x \in \mathbb{R}$ la successione a_n tende ancora a zero, ma oscillando perché due termini consecutivi hanno segno opposto; l'ampiezza delle oscillazioni, definita come la differenza, $|a_n - a_{n-1}|$ diminuisce con velocità esponenziale. Infatti:

$$|a_n - a_{n-1}| = |a_{2p} - a_{2p-1}| = |\lambda|^{2p-1} |\lambda - 1| < 1$$

4. $\lambda < -1$. Qui abbiamo $a_{2p} = |\lambda|^{2p}x \rightarrow +\infty$ per $p \rightarrow \infty$ crescendo, mentre $a_{2p+1} = -|\lambda|^{2p+1}x \rightarrow -\infty$ decrescendo, cioè da valori negativi. Quindi $\forall x \in \mathbb{R}$ la successione a_n oscillando perché due termini consecutivi hanno segno opposto; l'ampiezza delle oscillazioni, definita come la differenza, $|a_n - a_{n-1}|$ aumenta con velocità esponenziale. Infatti:

$$|a_n - a_{n-1}| = |a_{2p} - a_{2p-1}| = |\lambda|^{2p-1}|\lambda - 1| > 1$$

In simili casi si dice che la successione è *divergente per oscillazione*.

5. Dato che $\lambda^{-n} = \frac{1}{\lambda^n}$, se scambiamo n con $-n$ i ruoli di λ e λ^{-1} si invertono, nel senso che le divergenze si trasformano in convergenze e viceversa. Poiché lo scambio $n \leftrightarrow -n$ corrisponde a scambiare il passato col futuro, possiamo concludere che l'inversione del tempo trasforma convergenze esponenziali in divergenze esponenziali e viceversa.

Dunque in questo esempio siamo in grado di classificare il comportamento di *tutti* i dati iniziali $x \in \mathbb{R}$ per $n \rightarrow \infty$ (e anche per $n \rightarrow -\infty$ in virtù dell'osservazione precedente). Oltre i casi $\lambda = \pm 1$ già classificati, si vede che ogni dato iniziale converge verso il punto fisso $x = 0$ con velocità esponenziale per $n \rightarrow +\infty$ se $|\lambda| < 1$, e se ne allontana con pari velocità se $|\lambda| > 1$. Per $n \rightarrow -\infty$, cioè nel passato, questi comportamenti si invertono. In un linguaggio che formalizzeremo meglio fra poco, il punto fisso a $x = 0$ è *stabile e globalmente attrattivo* (nel futuro, o in avanti) se $|\lambda| < 1$, e *instabile e globalmente repulsivo* se $|\lambda| > 1$.

Il significato del concetto di stabilità di un punto fisso è intuitivo: punti che si scostano poco da $x = 0$ per $n = 0$, cioè che all'istante iniziale si scostano poco dal punto fisso, continuano a scostarsene poco per tutti gli istanti successivi; *attrattivo* significa che i dati iniziali che si scostano poco da $x = 0$ non solo gli rimangono vicini per ogni n , ma che convergeranno a tale punto per $n \rightarrow \infty$; *globalmente attrattivo* significa che *ogni* dato iniziale convergerà a $x = 0$, anche quelli inizialmente distanti a piacere da tale punto. Analogamente, *instabilità* significa che un dato iniziale vicino a $x = 0$ se ne allontana al crescere di n , e *repulsività* significa che esso tende ad allontanarsene indefinitamente. Repulsività globale significa che tutti i dati iniziali si comportano così.

Esempio 5.24

Consideriamo ora la mappa logistica $F_\mu(x) = \mu x(1-x)$ della (1.5), con $\mu > 1$. Sappiamo che essa ammette i punti fissi $x = 0$ e $x = p_\mu := \frac{\mu-1}{\mu}$. Facciamo vedere quanto segue:

$$\lim_{n \rightarrow \infty} F_\mu^n(x) \rightarrow -\infty \quad \text{se } x < 0 \quad \text{oppure } x > 1.$$

Infatti se, per cominciare, $x < 0$ si ha $\mu x(1-x) < x$ perché $1-x > 1$ e quindi $F_\mu(x) < x$. Poiché questa affermazione vale per ogni $x < 0$ possiamo applicarla di nuovo e concludere che la successione $F_\mu^n(x)$ è decrescente. Ora non esiste alcun $p \in \mathbb{R}$ per cui $F_\mu^n(p)$ converge a p . Se così fosse, infatti, si avrebbe anche $\lim_{n \rightarrow +\infty} F_\mu^{n+1}(p) = F_\mu(p) < p$, in contraddizione con l'affermazione precedente. Pertanto per il Teorema 3.2 la successione $F_\mu^n(p)$ tende a $-\infty$ per ogni $p < 0$. Se poi $x > 1$ si ha ancora $F_\mu(x) < x$ e quindi si può ripetere lo stesso ragionamento. La conclusione è che tutte le orbite corrispondenti ai dati iniziali $x < 0$ e $x > 1$ sono classificate, perché tendono tutte a $-\infty$. In particolare possiamo concludere che il punto fisso a $x = 0$ non è stabile.

Esempio 5.25

Consideriamo la dinamica discreta definita dalla ricorrenza di Fibonacci $a_{n+2} = a_{n+1} + a_n$ con le condizioni iniziali $a_0 = p, a_1 = q, (p, q) \in \mathbb{R}^2$. La forma esplicita di ogni orbita $a_n(p, q)$ è data dalla (4.17) con $a = b = 1$, che qui riscriviamo per comodità:

$$A_n = \frac{q - p(1 - \sqrt{5})/2}{\sqrt{5}} \left[\frac{1 + \sqrt{5}}{2} \right]^n + \left(p - \frac{q - p(1 - \sqrt{5})/2}{\sqrt{5}} \right) \left[\frac{1 - \sqrt{5}}{2} \right]^n \quad (5.10)$$

Si tratta di una combinazione lineare di due leggi esponenziali, l'uno crescente, perché $\lambda_1 := \frac{1 + \sqrt{5}}{2} > 1$, e l'altro decrescente, perché $|\lambda_2| = \left| \frac{1 - \sqrt{5}}{2} \right| < 1$; pertanto il punto fisso $(p, q) = (0, 0)$ non può essere stabile. Per $\epsilon > 0$ arbitrario, possiamo infatti trovare infiniti punti nel cerchio $p^2 + q^2 \leq \epsilon$ per cui il coefficiente di λ_1^n non è zero e quindi l'orbita si allontana indefinitamente dall'origine.

Studiamo tuttavia in maggior dettaglio la dipendenza dai dati iniziali. I termini esponenzialmente crescente e decrescente sono rispettivamente nulli per

$$q - p(1 - \sqrt{5})/2 = 0, \quad -q + p(1 + \sqrt{5})/2 = 0$$

Abbiamo già osservato che si tratta di due rette ortogonali passanti per l'origine nel piano (p, q) . Se il dato iniziale si trova sulla prima, esso convergerà con velocità esponenziale al

punto fisso nell'origine $p = q = 0$ per $n \rightarrow +\infty$; se si trova sulla seconda, se ne allontanerà esponenzialmente. Un punto fisso che ammette due direzioni ortogonali che ivi si incrociano, su una delle quali la dinamica evolve in maniera attrattiva e sull'altra in maniera repulsiva, si dice iperbolico.

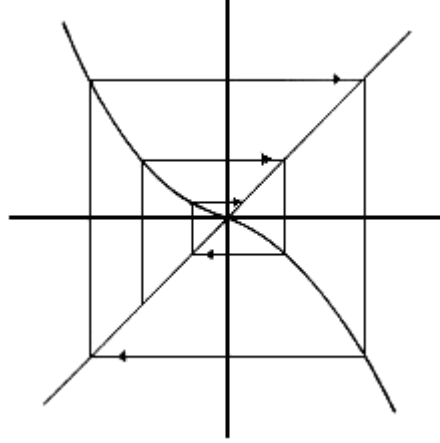


Figure 3.1: Esempio di punto fisso iperbolico.

Si noti che i ruoli delle due rette si invertono per $n \rightarrow \infty$: stabilità nel "futuro" ($n > 0$) diventa instabilità nel "passato" ($n < 0$) e viceversa.

Gli esempi precedenti motivano la seguente definizione generale, di cui ci serviremo spesso in seguito:

Definizione 5.18 Sia $y = (y_1, \dots, y_p) \in \mathbb{R}^p$ un punto fisso per la dinamica discreta definita dalla (3.8). Allora

1. Il punto fisso x si dice stabile se, $\forall \epsilon > 0$, esiste $\delta(\epsilon)$ tale che $|F^n(x) - y| \leq \epsilon$ per tutti gli x appartenenti alla sfera $S_\epsilon(y)$ di centro y e raggio $\delta(\epsilon)$, cioè tutti gli x tali che $\|x - y\| := \sqrt{(x_1 - y_1)^2 + \dots + (x_p - y_p)^2} < \delta(\epsilon)$.
2. Il punto fisso stabile y si dice poi attrattivo, o attrattore se

$$\lim_{n \rightarrow +\infty} |F^n(x) - y| = 0, \quad \forall x \in S_\epsilon(y) \quad (5.11)$$

In tal caso, ogni punto $x \in S_\epsilon(y)$ si dice asintotico (nel futuro) al punto fisso y .

3. Se la stabilità vale per ogni x del dominio di definizione di F il punto fisso si dice globalmente stabile; se vale inoltre la (5.11) il punto fisso si dice globalmente attrattivo o attrattore globale.

4. Se per ogni punto $x \in S_\epsilon(y)$ si ha

$$\lim_{n \rightarrow +\infty} |F^n(x) - y| = +\infty \quad (5.12)$$

il punto fisso si dice repulsivo o repulsore; se la (5.12) vale per ogni x del dominio di F il punto fisso si dice globalmente repulsivo o repulsore globale.

5. L'insieme dei punti asintotici nel futuro si dice insieme stabile del punto fisso p . L'insieme dei punti asintotici nel passato si dice insieme instabile del punto fisso p .

Osservazione 5.27

Il dato iniziale y è un punto di $\mathbb{R}^p$: $y = (y_1, \dots, y_p)$. La distanza $|F^n(x) - y|$ è, per definizione, la distanza massima fra le componenti cartesiane di y e il numero reale $F^n(x)$:

$$|F^n(x) - y| := \max_{1 \leq k \leq p} |F^n(x) - y_k|$$

Esempio 5.26

Esempi di punti fissi attrattivi e repulsivi ne abbiamo già visti. Vediamo ora un esempio di punto fisso stabile ma non attrattivo. Consideriamo ancora la dinamica discreta definita dalla ricorrenza $a_{n+2} + a_n = 0$ del §?, per cui ogni dato iniziale $\alpha = (\alpha_1, \alpha_2) \in \mathbb{R}^2$ è periodico come abbiamo visto in precedenza e che ammette il punto fisso $(0, 0)$. Facciamo vedere che questo punto fisso è stabile ma non attrattivo. Infatti, fissati i dati iniziali $a_0 = \alpha_0$; $a_1 = \alpha_1$ l'orbita è $a_{2n} = (-1)^n \alpha_0$, $a_{2n+1} = (-1)^n \alpha_1$. La condizione $|F^n(x) - y| < \epsilon$ si scrive qui

$$\sqrt{|(-1)^n \alpha_0|^2 + |(-1)^n \alpha_1|^2} = \sqrt{\alpha_0^2 + \alpha_1^2} < \epsilon;$$

ciò è sicuramente vero per $\sqrt{\alpha_0^2 + \alpha_1^2} < \delta(\epsilon)$ con $\delta(\epsilon) = \epsilon$, cioè per (α_0, α_1) nel cerchio di centro l'origine e raggio ϵ (sfera nel piano). Il punto fisso non è però attrattivo perché, come abbiamo già visto, tutte le successioni $a_n(\alpha)$ sono periodiche; pertanto non hanno limite e non possono convergere a 0 per $n \rightarrow \infty$.

Osservazione 5.28

Consideriamo ancora l'esempio precedente. Si noti che l'insieme dei dati iniziali tali che $\sqrt{\alpha_0^2 + \alpha_1^2} < \epsilon$ è un cerchio di raggio $\sqrt{\epsilon}$ e quindi di area $\pi\epsilon$. Calcoliamo ora l'immagine di questo cerchio attraverso la dinamica discreta definita dalla ricorrenza $a_{n+2} + a_n = 0$ ad ogni "istante" n , e facciamo vedere che si tratta ancora del medesimo cerchio, che quindi avrà la medesima area. Notiamo così che la dinamica discreta conserva le aree.

All'istante $n, n = 1, 3, \dots$ l'immagine della coppia di dati iniziali α_0 e α_1 è la coppia $a_{2n} = (-1)^n \alpha_0, a_{2n+1} = (-1)^n \alpha_1$. Pertanto

$$\sqrt{|a_{2n}|^2 + |a_{2n+1}|^2} = \sqrt{\alpha_0^2 + \alpha_1^2}$$

e quindi l'area di ogni cerchio è conservata.

Capitolo 4

Studio della mappa logistica: frattali e caos

In questo capitolo esamineremo molto più in dettaglio la dinamica discreta generata dalla mappa logistica. La ragione è che questo modello semplicissimo da formulare permette di dare un esempio esplicito e sorprendentemente generale di strutture matematiche assai profonde e importanti quali i frattali ed il comportamento caotico delle orbite.

1 L'iperbolicità dei punti periodici

Abbiamo già visto un esempio di punto fisso iperbolico. Vediamo ora di specializzare questo concetto al caso dei punti periodici della dinamica discreta unidimensionale definita da una funzione $f(x) : I \subset \mathbb{R} \rightarrow \mathbb{R}$, che ammetteremo sempre derivabile tante volte quante si vuole su tutto l'intervallo I .

Definizione 1.19 *Sia p un punto periodico per f di periodo primitivo N : $f^N(p) = p$. Il punto p si dice iperbolico se $|(f^N)'(p)| \neq 1$. (Qui l'apice denota la derivata).*

Osservazione 1.29

1. Per $N = 1$ si ha un punto fisso, che sarà quindi iperbolico se $|f'(p)| \neq 1$.
2. Cerchiamo di rendere esplicito il significato della condizione di iperbolicità. Poniamo per comodità di notazione $f^N(x) = g(x)$. Allora p è un punto fisso di g : $g(p) = p$. Per $|x - p|$ sufficientemente piccolo applichiamo la formula di Taylor di punto iniziale p al secondo ordine. Ricordiamo dall'analisi la formula di Lagrange-Taylor all'ordine n : se

$f(x)$ è derivabile almeno $n + 1$ volte in un intervallo aperto $I \subset \mathbb{R}$, e $x - x_0 \in I$, esiste un intervallo $J(x_0) \subset I$ tale che per $x \in J$ si può scrivere:

$$\begin{aligned} f(x) &= f(x_0) + f'(x_0)(x - x_0) + f''(x_0)\frac{(x - x_0)^2}{2} + \dots + f^{(n)}(x_0)\frac{(x - x_0)^n}{n!} \\ &+ \frac{f^{(n+1)}(\bar{x})(x - x_0)^{n+1}}{(n+1)!} \end{aligned}$$

dove il punto $\bar{x} \in I$, che dipende da x e x_0 , è tale che $0 < |\bar{x}| < |x - x_0|$. Si ottiene allora, se si tronca la formula tenendo il resto al secondo ordine:

$$g(x) = g(p) + g'(p)(x - p) + g''(\bar{x})(x - p)^2 = [g(p) - pg'(p)] + g'(p)x + g''(\bar{x})(x - p)^2$$

Traslando l'origine di $g(p) - pg'(p)$ sull'asse delle ordinate e trascurando i termini quadratici otteniamo per definizione la *linearizzazione*, o *parte lineare*, o *mappa linearizzata* della mappa definita dalla funzione $f^N(x)$ attorno al punto fisso p :

$$f_L(p) := g'(p)x \quad (1.1)$$

Allora la condizione di iperbolicità significa che la mappa linearizzata è una mappa non banale, cioè diversa $\pm x$ (x è la mappa identica, e $-x$ la mappa identica cambiata di segno). Ad esempio, la funzione $f(x) = \sin x$ genera una mappa non lineare su $[0, 2\pi]$ i cui punti fissi a $x = 0$ e $x = \pi$ non sono iperbolicici perché $f'(0) = \cos(0) = 1$, $f'(\pi) = \cos(\pi) = -1$. La mappa logistica genera invece punti fissi iperbolicici come ora vedremo.

Esempio 1.27

1. Consideriamo la mappa logistica $f_\mu(x) = \mu x(1 - x)$ per $\mu > 1$. Ci sono come sappiamo due punti fissi: $x = 0$ e $x = p_\mu = \frac{\mu - 1}{\mu}$. Allora se $\mu \neq 2$ entrambi sono iperbolicici perché $f'_\mu(0) = \mu$ e $f'_\mu(p_\mu) = 2 - \mu$;
2. Sia $f(x) = -\frac{1}{2}(x^3 + x)$. Qui c'è un punto fisso a $x = 0$, iperbolicico perché $f'(0) = -\frac{1}{2}$, e due punti periodici di periodo 2, $x = \pm 1$. Questi sono entrambi iperbolicici, perché $(f \circ f)'(\pm 1) = f'(1) \cdot f'(-1) = 4$.

Abbiamo già completamente analizzato il comportamento della mappa linearizzata f_L per quanto riguarda l'analisi di stabilità: attrattività globale del punto fisso $x = 0$ se $|g'(p)| < 1$,

repulsività globale se $|g'(p)| > 1$. Poiché la nostra mappa f è approssimabile con la mappa lineare f_L con approssimazione tanto migliore quanto più x si avvicina al punto fisso p , ci si aspetta in corrispondenza dei due casi $|g'(p)| < 1$ e $|g'(p)| > 1$ rispettivamente l'attrattività e la repulsività *locale*. Proprietà *locale* significa che deve esistere almeno un intorno del punto fisso i cui punti godono della proprietà in questione. Per *intorno* di un punto in $x \in \mathbb{R}$ intenderemo qui un intervallo aperto centrato in x . Nel seguito quando parleremo di proprietà quali l'attrattività ecc. senza alcuna ulteriore specificazione esse saranno sempre intese in senso locale. Qui si ha:

Proposizione 1.5

1. Sia p un punto fisso iperbolico di f con $|f'(p)| < 1$. Allora p è stabile e attrattivo.
2. Sia p un punto fisso iperbolico di f con $|f'(p)| > 1$. Allora p è instabile e repulsivo.

Dimostrazione. Dobbiamo far vedere che esiste un intervallo aperto I attorno a p tale che $\lim_{n \rightarrow +\infty} f^n(x) = p \forall x \in I$. Poiché f è almeno C^1 (cioè ammette derivata prima continua sul suo dominio di definizione), $\exists \epsilon > 0$ tale che $|f'(x)| < A < 1$ per $x \in [p - \epsilon, p + \epsilon]$. Ora applichiamo il teorema del valor medio di Lagrange tenendo conto del fatto che p è un punto fisso di f . Si ottiene:

$$|f(x) - p| = |f(x) - f(p)| \leq A|x - p| < |x - p| \leq \epsilon$$

Pertanto $f(x) \in [p - \epsilon, p + \epsilon]$ e quindi la distanza di $f(x)$ da p è minore della distanza di x stesso da p . Iterando il ragionamento troviamo

$$|f^n(x) - p| \leq A^n|x - p|$$

e pertanto $\lim_{n \rightarrow +\infty} f^n(x) = p$ perchè $A < 1$.

Se $|f'(p)| > 1 \exists \epsilon > 0$ tale che $|f'(x)| > A > 1$ per $x \in [p - \epsilon, p + \epsilon]$. Ripetendo il ragionamento otteniamo la disuguaglianza

$$|f^n(x) - p| \geq A^n|x - p|$$

Dunque la distanza fra $f^n(x)$ e p tende all'infinito. Ciò condude la prova della proposizione.

Osservazione 1.30

Un punto periodico di periodo $N \geq 2$ di $f(x)$ è per definizione un punto fisso della mappa $g(x) := f^N(x)$. A quest'ultimo potremo quindi applicare la proposizione precedente. È quindi naturale porre la seguente ulteriore

Definizione 1.20 *Sia p un punto periodico iperbolico di periodo primitivo N della mappa f , e sia $g = f^N$. Allora*

1. *Diremo che p è stabile ed attrattivo se $|g'(p)| = |(f^N)'(p)| < 1$.*
2. *Diremo che p è instabile e repulsivo se $|g'(p)| = |(f^N)'(p)| > 1$.*
3. *Gli insiemi stabili ed instabili del punto periodico sono gli insiemi stabili ed instabili del punto fisso di g .*

Per definizione un punto periodico p (e in particolare fisso) per la mappa f sarà non iperbolico quando $f'(p) = \pm 1$.

In mappe che, come quella logistica, dipendono da un parametro, può succedere che al variare del parametro medesimo un punto fisso non iperbolico possa dare luogo ad una coppia di punti fissi iperbolici. Questo fenomeno, quando avviene, si chiama *biforcazione*.

Poiché darne una definizione generale può essere inutilmente complicato, ci limiteremo ad illustrare questa nozione tramite due esempi:

Esempio 1.28

1. Consideriamo la famiglia di mappe quadratiche $q_c(x) := x^2 + c$, $c \in \mathbb{R}$. Se $c > \frac{1}{4}$ q_c non ammette punti fissi perché l'equazione $x = x^2 + c$ non ha radici reali: infatti, dette $x_{1,2}$ le radici, si ha $x_{1,2} = \frac{1}{2}[1 \pm \sqrt{1 - 4c}]$. Se $c = \frac{1}{4}$ nasce un punto fisso a $x = \frac{1}{2}$, che non è iperbolico perché $q'_c(\frac{1}{2}) = 1$. Per $c < \frac{1}{4}$ vi sono due punti fissi; poiché $q'_c(x_{1,2}) = [1 \pm \sqrt{1 - 4c}]$, x_1 è repulsivo e x_2 è attrattivo. Si ha quindi una biforcazione a $c = \frac{1}{4}$.
2. Consideriamo ancora la mappa logistica $f_\mu(x) = \mu x(1 - x)$ con $\mu > 1$ ed i suoi punti fissi 0 e $p_\mu = \frac{\mu - 1}{\mu}$. Ora $f'_\mu(0) = \mu$, $f'_\mu(p_\mu) = 2 - \mu$. Quindi 0 è sempre repulsivo ($\mu > 1$) mentre p_μ è attrattivo per $1 < \mu < 3$. Quando $\mu = 3$ si ha $f'_\mu(p_\mu) = -1$ ed il punto fisso p_μ perde il carattere di iperbolicità. Consideriamo ora la mappa $f_\mu^2(x)$,

cioè μ composta con sè stessa. Si può dimostrare che se $\mu > 3$ essa ha due punti fissi iperbolici, che quindi saranno punti periodici iperbolici di f_μ di periodo 2, uno attrattivo e l'altro repulsivo. Questo è un altro esempio di biforcazione.

Esercizio 1.27

Dimostrare l'affermazione precedente.

(Suggerimento: dimostrare che

$$g_\mu(x) := f_\mu^2(x) = \mu^2 x(1-x)(1-\mu x(1-x))$$

e fare vedere che per $\mu = 3$ la retta $y = x$ è tangente alla curva $y = g_\mu(x)$).

2 La mappa logistica al crescere del parametro

Cominciamo in questo paragrafo lo studio ancor più dettagliato della mappa logistica $f_\mu(x) = \mu x(1-x)$ per $\mu > 1$. Abbiamo già dimostrato che:

1. $\lim_{n \rightarrow +\infty} f_\mu^n(x) = -\infty$ se $x < 0$ o $x > 1$.
2. Se $\mu < 3$ f_μ ha un punto fisso repulsivo a $x = 0$ e un punto fisso attrattivo a $x = p_\mu := \frac{\mu-1}{\mu}$.

L'affermazione 1 mostra che le sole condizioni iniziali interessanti sono quelle per cui $0 \leq x \leq 1$. Questo fatto è perfettamente in linea con l'interpretazione di x come popolazione relativa. 0 è un punto fisso. La proposizione che segue classifica il destino di ogni condizione iniziale quando $\mu < 3$, cioè prima della prima biforcazione.

Proposizione 2.6 *Sia $0 < x < 1$. Allora:*

$$\lim_{n \rightarrow +\infty} f_\mu^n(x) = p_\mu \tag{2.2}$$

Osservazione 2.31

Dunque, a parte il punto fisso a $x = 0$ (soluzione banale), ogni dato iniziale della ricorrenza a due termini tende a p_μ per $n \rightarrow \infty$. Se si prende sul serio la mappa come modello qualitativo della crescita delle popolazioni (relative), questa proposizione afferma che per tutti i valori

del parametro fra 1 e 3 ci sarà un solo valore di equilibrio per la popolazione relativa a cui si giungerà quale che fosse la popolazione relativa iniziale.

Dimostrazione. Cominciamo col supporre $1 < \mu < 2$, e $x \in]0, 1/2]$. Si noti che $1/2 < p_\mu < 1$ e che per $x \in]0, 1/2]$ la $f_\mu(x)$ è iniettiva. Infatti $f'_\mu(x) = \mu(1 - 2x) > 0$ per $x \in]0, 1/2[$. Pertanto $f_\mu(x)$ è strettamente crescente e come tale iniettiva. Allora in tale intervallo si ha

$$|f_\mu(x) - p_\mu| < |x - p_\mu|, \quad x \neq p_\mu$$

Distinguiamo i 4 casi che occorre discutere per risolvere la disuguaglianza precedente:

1. $f_\mu(x) - p_\mu > 0$, $x - p_\mu > 0$. Poiché per $x \in]0, 1/2[$ e $1 < \mu < 2$ si ha $x > \mu x(1 - x) = f_\mu(x)$; quindi $f_\mu(x) - p_\mu < x - p_\mu$.
2. $f_\mu(x) - p_\mu < 0$, $x - p_\mu > 0$. Risolvendo la disequazione di secondo grado $\mu x(1 - x) = p_\mu < 0$ si trova che $f_\mu(x) - p_\mu < 0$ per $x < p_\mu$ oppure $x > 1/\mu$. Poiché $1/\mu > 1/2$ e $x > p_\mu$ questo caso non si verifica mai.
3. $f_\mu(x) - p_\mu > 0$, $x - p_\mu < 0$. La seconda disuguaglianza implica $x < p_\mu$, e per questi valori di x sappiamo dal caso precedente che $f_\mu(x) - p_\mu < 0$. Quindi nemmeno il caso 3 si verifica.
4. Questo caso è equivalente al caso 1.

Dunque la disuguaglianza è dimostrata. Allora, iterando, tramite un'ulteriore applicazione del ragionamento impiegato già molte volte, possiamo concludere che $\lim_{n \rightarrow +\infty} f_\mu^n(x) = p_\mu$. D'altra parte, se $x \in]1/2, 1[$ allora $0 < f_\mu(x) < 1/2$ cosicché il ragionamento precedente implica

$$f_\mu^n(x) = (f_\mu^{n-1} \circ f_\mu)(x) \rightarrow p_\mu$$

Il caso $2 < \mu < 3$ è più complicato. Denotiamo $\tilde{p}_\mu$ l'unico punto nell'intervallo $]0, 1/2[$ la cui immagine rispetto a $f_\mu(x)$ è p_μ . Si può verificare tramite un calcolo esplicito, che ometteremo, che l'immagine dell'intervallo $[\tilde{p}_\mu, p_\mu]$ attraverso $f_\mu \circ f_\mu$ è contenuta nell'intervallo $[1/2, p_\mu]$. Ne segue, come sopra, che $\lim_{n \rightarrow +\infty} f_\mu^n(x) = p_\mu \forall x \in [\tilde{p}_\mu, p_\mu]$. Supponiamo ora $x < \tilde{p}_\mu$. Si può dimostrare, con un ulteriore calcolo esplicito che ometteremo ancora, che esiste $k > 0$ tale che $f^k(x) \in [\tilde{p}_\mu, p_\mu]$. Dunque anche in questo caso $\lim_{n \rightarrow +\infty} f^{k+n}(x) = p_\mu$. Finalmente; come sopra, l'immagine di $]p_\mu, 1[$ attraverso f_μ è esattamente l'intervallo $]0, p_\mu]$. Quindi il risultato vale anche per $x \in]p_\mu, 1[$. Ora notiamo che

$$]0, 1[=]0, \tilde{p}_\mu[\cup [\tilde{p}_\mu, p_\mu] \cup]p_\mu, 1[$$

e questo conclude la dimostrazione, a parte il caso del punto $\mu = 2$ che viene lasciato come esercizio.

Esercizio 2.28

Sia $f_2(x) = 2x(1-x)$. Dimostrare che $\lim_{n \rightarrow +\infty} f_2^n(x) = \frac{1}{2}$ per ogni $0 < x < 1$.

Soluzione. $x = \frac{1}{2}$ è un punto fisso. Sia $\epsilon > 0$. Facciamo vedere anzitutto che l'immagine dell'intervallo $\left] \epsilon, \frac{1}{2} \right]$ è un intervallo $\left] \epsilon_1, \frac{1}{2} \right]$ strettamente contenuto in $\left] \epsilon, \frac{1}{2} \right]$. A questo scopo, dato che $\frac{1}{2}$ è un punto fisso basta far vedere che l'immagine ϵ_1 di ϵ attraverso f_2 è tale che $\epsilon < \epsilon_1 < \frac{1}{2}$. Si ha: $f_2(\epsilon) = 2\epsilon(1-\epsilon) := \epsilon_1$; $\epsilon_1 > \epsilon$ dato che $0 < \epsilon < \frac{1}{2}$. Iterando il ragionamento, al secondo passo troviamo un intervallo $\left] \epsilon_2, \frac{1}{2} \right]$ con $\epsilon_2 > \epsilon_1$ e per $n \rightarrow \infty$ concludiamo che ogni dato iniziale $0 < x \leq \frac{1}{2}$ convergerà a 1 per $n \rightarrow +\infty$. Se $x \in \left[\frac{1}{2}, 1 - \epsilon \right]$ vale lo stesso ragionamento perché $f_2(x)$ è simmetrica rispetto alla sostituzione $x \leftrightarrow 1-x$ essendo $\frac{1}{2}$ il centro dell'intervallo $[0, 1]$.

Per $\mu = 3$ sappiamo che ha luogo una biforcazione. C'è quindi da attendersi un comportamento più complicato per $\mu > 3$. Questo è esattamente quello che avviene. Per descriverlo, è bene premettere una nozione importante.

3 L'insieme ternario di Cantor

Procediamo alla costruzione dell'insieme di Cantor¹ Sia I l'intervallo chiuso $[0, 1] := \{x \in \mathbb{R} \mid 0 \leq x \leq 1\}$. Il procedimento iterativo procede nel modo seguente:

1. Togliamo da I l'intervallo aperto centrale $A_1^{(0)} := \left] \frac{1}{3}, \frac{2}{3} \right[$ di centro $\frac{1}{2}$ e lunghezza $\frac{1}{3}$ (la "terza parte centrale" aperta). (Usiamo la locuzione "togliere da" come abbreviazione di "considerare l'insieme complementare di". In altre parole, qui dovremmo dire: consideriamo il complementare rispetto ad I dell'intervallo aperto $\left] \frac{1}{3}, \frac{2}{3} \right[$). Rimane l'insieme complementare $I - A_1^{(0)}$ definito dall'unione disgiunta di intervalli chiusi $\left[0, \frac{1}{3} \right] \cup \left[\frac{2}{3}, 1 \right]$.

¹Georg Cantor, (S.Pietroburgo 1845- Halle a.d. Saale (Germania) 1918), Professore di Matematica all'Università di Halle-Wittenberg, fu il fondatore della moderna teoria degli insiemi. A lui si debbono le nozioni di potenza di un insieme, numero transfinito, ecc. L'insieme ternario che ora procederemo a costruire è una delle nozioni più sottili e profonde della matematica.

2. Dall'insieme rimasto (il complementare $I - A_1^{(0)}$) togliamo ora le due terze parti centrali aperte, e cioè gli intervalli aperti $]\frac{1}{9}, \frac{2}{9}[:= A_1^{(1)}$ e $]\frac{7}{9}, \frac{8}{9}[:= A_2^{(1)}$ di lunghezza $\frac{1}{9}$ e centro $\frac{1}{6}$ e $\frac{5}{6}$, rispettivamente. Rimane l'insieme complementare $I - A_1^{(0)} - A_1^{(1)} - A_2^{(1)}$ definito dall'unione disgiunta di intervalli chiusi $[0, \frac{1}{9}] \cup [\frac{2}{9}, \frac{1}{3}] \cup [\frac{7}{9}, 1] \cup [\frac{8}{9}, 1]$.
3. Ripetiamo indefinitamente il procedimento.

Si ha allora

Lemma 3.1

1. Al passo di ordine n dell'iterazione, $n = 0, 1, 2, \dots$ si rimuovono esattamente 2^n intervalli aperti $A_k^{(n)}$, $k = 1, \dots, 2^n$ ciascuno di lunghezza $\frac{1}{3^{n+1}}$.
2. Tutti gli intervalli $A_k^{(n)}$ sono disgiunti due a due: $A_k^{(n)} \cap A_l^{(m)} = \emptyset$, $k \neq l, n \neq m$.
3. Sia:

$$A_n := \bigcup_{k=1}^{2^n} A_k^{(n)}; \quad A := \bigcup_{n=0}^{\infty} A_n \quad (3.3)$$

Sia $l(A_n)$ la lunghezza totale di questa unione di intervalli disgiunti. Allora

$$l(A_n) = \frac{1}{3} \left(\frac{2}{3}\right)^n; \quad l(A) = \sum_{n=0}^{\infty} l(A_n) = 1$$

Dimostrazione. Dimostriamo la prima affermazione per induzione. Essa è vera per $n = 0$. Assumiamola vera per $s = 1, \dots, n$ e proviamola per $s = n + 1$. La lunghezza della parte centrale di ogni intervallo $A_k(n + 1)$ all'ordine $n + 1$ vale $\frac{1}{3}$ della lunghezza di ogni intervallo $A_k(n)$, e quindi vale $\frac{1}{3^{n+1}}$ perché la lunghezza di ogni intervallo all'ordine n è 3^{-n-1} . Gli intervalli $A_k(n)$ sono poi ad ogni ordine disgiunti per costruzione. Per provare la terza affermazione, per prima cosa notiamo che la lunghezza dell'unione di un numero qualsiasi di intervalli *disgiunti* è la somma delle loro lunghezze. Ad esempio, la lunghezza dell'unione dei due intervalli $] - 1, 0[$ e $]0, 1[$ vale $1 + 1 = 2$. La lunghezza dell'unione dei due intervalli $] - \frac{1}{2}, \frac{1}{2}[$ e $]0, 1[$, che non sono disgiunti, vale $\frac{3}{2} < 2$. In generale si può dimostrare che se A è un'unione finita o infinita di intervalli disgiunti I_k , cioè $A = \bigcup_{k=0}^{\infty} I_k$, e A è un insieme

limitato, allora la serie $\sum_{k=0}^{\infty} l(I_k)$ converge e si ha

$$l(A) = \sum_{k=0}^{\infty} l(I_k).$$

Pertanto:

$$l(A_n) = l\left(\sum_{k=1}^{2^n} l(A_k^{(n)})\right) = \frac{2^n}{3^{n+1}}$$

perché ognuno degli intervalli $A_k^{(n)}$ è lungo 3^{-n-1} , e ce ne sono 2^n . Dunque, poiché anche gli insiemi A_n sono disgiunti a due a due:

$$\begin{aligned} l(A) &= \sum_{n=0}^{\infty} l(A_n) = \sum_{k=0}^{\infty} \frac{2^k}{3^{k+1}} \\ &= \frac{1}{3} \sum_{k=0}^{\infty} \frac{2^k}{3^k} = \frac{1}{3} \frac{1}{1 - 2/3} = 1 \end{aligned}$$

e ciò conclude la prova del Lemma.

Siamo ora in grado di definire l'insieme ternario di Cantor.

Definizione 3.21 Il complementare rispetto a $I = [0, 1]$ dell'insieme $A := \bigcup_{n=1}^{\infty} A_n$ definito dalla (3.3) si dice insieme ternario di Cantor $\mathcal{C}$.

Osservazione 3.32

1. Per costruzione l'insieme di Cantor è un sottoinsieme di $[0, 1]$ chiuso, e *totalmente disconnesso*, cioè non contiene alcun intervallo. $\mathcal{C}$ è chiuso perché è il complementare dell'insieme aperto $\mathcal{A}$ rispetto all'intervallo chiuso $[0, 1]$. Infatti il complementare di un insieme aperto rispetto ad un insieme chiuso è chiuso, e a sua volta $\mathcal{A}$ è aperto perché si può dimostrare che l'unione finita o numerabile di intervalli aperti è un insieme aperto. (Per il significato della locuzione numerabile si veda il punto 3 successivo). Inoltre si può dimostrare un'altra proprietà che emerge dal meccanismo di costruzione, e cioè che si tratta di un insieme *perfetto*. Questo significa che ogni punto dell'insieme di Cantor è punto limite di successioni di punti appartenenti punti all'insieme medesimo.
2. Supponiamo di avere due sottoinsiemi B_1 e B_2 dell'intervallo $[0, 1]$ complementari rispetto all'intervallo: $B_1 \cap B_2 = \emptyset$, $B_1 \cup B_2 = [0, 1]$. Allora da quanto precede segue

che $l(B_1) + l(B_2) = l([0, 1]) = 1$. Ne concludiamo che la lunghezza dell'insieme di Cantor è zero poiché la lunghezza del suo insieme complementare rispetto all'intervallo $[0, 1]$ vale 1.

3. Tuttavia l'insieme di Cantor, anche se di lunghezza 0, "contiene lo stesso numero di punti" dell'intervallo $[0, 1]$. Si può infatti dimostrare che l'insieme di Cantor e l'intervallo $[0, 1]$ sono insiemi *equipotenti*, cioè che i punti dell'uno possono essere messi in corrispondenza biunivoca con quelli dell'altro. *La proprietà di ammettere sottoinsiemi propri equipotenti è caratteristica di un insieme infinito, ed è spesso anzi assunta come definizione di insieme infinito medesimo.* Infatti è del tutto evidente che gli elementi di un qualsiasi insieme composto da un numero finito di elementi non potranno mai essere messi in corrispondenza biunivoca con gli elementi di un suo sottoinsieme *proprio*.

Facciamo invece vedere, come semplice esempio, che l'insieme $\mathbb{N}$ dei numeri naturali è un insieme infinito. Esso è infatti equipotente al suo sottoinsieme proprio costituito dai numeri pari. Per vedere questa affermazione, basta fare corrispondere 2 a 1, 4 a 2, e in generale $2n$ a n , cioè ad ogni numero il suo doppio. La corrispondenza è chiaramente biunivoca. Allo stesso modo si prova che $\mathbb{N}$ è equipotente al sottoinsieme proprio costituito dai numeri dispari. Si può dimostrare che $\mathbb{N}$ e $\mathbb{Q}$ sono equipotenti. Gli insiemi che hanno la potenza dell'insieme $\mathbb{N}$ si dicono *numerabili*. L'intervallo $[0, 1]$ ha potenza strettamente maggiore del numerabile. Tutti gli insiemi equipotenti a $[0, 1]$ hanno per definizione la *potenza del continuo*. Dunque l'insieme di Cantor ha la potenza del continuo.

4. L'insieme di Cantor costituisce anche l'esempio più semplice di *insieme frattale*. Intuitivamente un insieme frattale è un insieme che rimane lo stesso se osservato in qualsiasi scala, o, più o meno equivalentemente, rimane lo stesso dopo qualsiasi ingrandimento, o anche qualsiasi suo sottoinsieme proprio ha la medesima struttura² Supponiamo ad esempio di osservare solo i punti che stanno nell'intervallo di sinistra $[0, \frac{1}{3}]$, ma attraverso un microscopio che ingrandisce da 1 a 3. Allora questo "pezzo" dell'insieme di Cantor è esattamente uguale all'insieme originale. Il procedimento può naturalmente

²In inglese questa proprietà si dice *self-similarity*.

essere iterato: allo stadio n ogni pezzo dell'insieme di Cantor è esattamente uguale all'insieme originale se lo si osserva con un microscopio che ingrandisce da 1 a 3^n .

Le proprietà precedenti, assunte in astratto, danno la definizione generale di insieme di Cantor:

Definizione 3.22 *Un sottoinsieme limitato Λ di $\mathbb{R}$ si dice insieme di Cantor se è chiuso, totalmente sconnesso e perfetto.*

La cosa sorprendente è che la costruzione di un insieme di Cantor, apparentemente frutto della pura invenzione matematica³, è generata dalle orbite della dinamica discreta anche nel suo esempio non lineare più semplice, la mappa logistica.

4 La mappa logistica genera un insieme di Cantor

Riconsideriamo la mappa logistica, stavolta con $\mu > 3$, e omettiamo di indicare la dipendenza esplicita da μ . Scriveremo dunque $f(x) = \mu x(1 - x)$, e sia $x \in I := [0, 1]$. Sappiamo infatti che se $x \notin [0, 1]$ ogni dato iniziale tende a $-\infty$. Cominciamo coll'osservare che, poiché $\mu > 4$, il massimo di f che vale $\frac{\mu}{4}$ è maggiore di 1. Dunque certi punti di I ne usciranno dopo un'iterazione. Denotiamo A_0 l'insieme di questi punti:

$$A_0 := \{x \in I \mid f(x) > 1\}$$

Ora il massimo di f è raggiunto nel punto $x = \frac{1}{2}$, e f è simmetrica attorno a tale punto. Pertanto, per continuità, A_0 sarà un intervallo aperto di centro $\frac{1}{2}$ contenuto in I . Se $x \in A_0$, si ha $f(x) > 1$ e quindi $(f \circ f)(x) < 0$ da cui, come sappiamo, $f^n(x) \rightarrow -\infty$ per $n \rightarrow \infty$. Dunque A_0 è l'insieme dei punti che escono subito (cioè alla prima applicazione della mappa f) da I per andarsene a $-\infty$ quando $n \rightarrow \infty$. Equivalentemente, tutti i punti in $I - A_0$ rimangono in I dopo la prima iterazione. Consideriamo ora l'insieme

$$A_1 := \{x \in I \mid f(x) \in A_0\}$$

³In realtà Cantor fu condotto ad isolare queste nozioni dalle ricerche che andava conducendo sulle proprietà degli insiemi di punti sui quali si cercava di simostrare la convergenza delle serie di Fourier delle funzioni continue.

Se $x \in A_1$ allora $(f \circ f)(x) > 1$, $(f \circ f \circ f)(x) < 0$ e quindi, come sopra, $f^n(x) \rightarrow -\infty$ per $n \rightarrow \infty$. Definiamo allora induttivamente la successione di insiemi

$$A_n := \{x \in I \mid f^n(x) \in A_0\} = \{x \in I \mid f^k(x) \in I, k = 1, \dots, n \mid f^{n+1}(x) \notin I\}$$

In altre parole, A_n consiste in tutti i punti che escono da I alla $n+1$ -esima iterazione. Come sopra, ne segue che l'orbita di x tenderà a $-\infty$ per $n \rightarrow \infty$. Poiché conosciamo quale sarà il destino dei punti che stanno in A_n , ci rimane da analizzare cosa capiterà a quelli che non escono mai da I ; vogliamo cioè analizzare l'orbita di tutti i dati iniziali che appartengono all'insieme

$$\Lambda := I - \left(\bigcup_{n=0}^{\infty} A_n \right) \quad (4.4)$$

È chiaro che la prima cosa da capire è come sia fatto questo insieme. Per farlo, guardiamo più da vicino la sua costruzione ricorsiva. A_0 è un intervallo aperto di centro $\frac{1}{2}$. Pertanto il suo complementare $I - A_0$ consiste di due intervalli chiusi disgiunti, uno a destra di $\frac{1}{2}$, denotato I_1 , e uno a sinistra, denotato I_0 . Notiamo poi che f è monotona tanto su I_0 (crescente) quanto su I_1 (decrescente), e che $f(I_0) = f(I_1) = I$. Infatti $f(I_0) = f(I_1)$ segue dalla simmetria tanto di f quanto di I_0 e I_1 attorno a $x = \frac{1}{2}$, e il fatto che l'immagine sia I è conseguenza della monotonia: $f(I_0)$ parte da 0 e arriva a 1 perché per definizione A_0 è l'insieme degli $x \in I$ per cui $f(x) > 1$. Ora dato che $f(I_0) = f(I_1) = I$ esiste una coppia di intervalli aperti, uno contenuto in I_0 e l'altro in I_1 , la cui immagine attraverso f è A_0 . Questa coppia di intervalli (disgiunti, e simmetrici rispetto a $x = \frac{1}{2}$) è per definizione l'insieme A_1 .

Consideriamo ora l'insieme complementare $I - (A_0 \cup A_1)$. Questo insieme consiste in 4 intervalli chiusi e f trasforma ciascuno di questi monotonicamente in I_0 oppure in I_1 . Dunque ragionando come prima, potremo affermare che $f^2 = f \circ f$ trasformerà ciascuno di questi in I . Ancora iterando il ragionamento precedente, vediamo che ciascuno dei 4 intervalli in $I - (A_0 \cup A_1)$ contiene un sottointervallo che viene trasformato in A_0 da f^2 . Pertanto i punti di questi intervalli escono da I dopo la terza iterazione, e formano per definizione l'insieme A_2 . Ci servirà in seguito l'osservazione che f^2 cresce e decresce alternativamente in ciascuno di questi 4 intervalli. Infatti, si ha

$$(f \circ f)' = f(f(x))f'(x) = \mu^2[1 - 2\mu x(1 - x)](1 - 2x)$$

L'ultimo fattore è positivo per $0 < x < 1/2$, e negativo per $1/2 < x < 1$; il fattore $1 - 2\mu x(1 - x)$ è negativo per $\frac{1}{2}(1 - \sqrt{1 - 2/\mu}) < x < \frac{1}{2}(1 + \sqrt{1 - 2/\mu})$. Quindi f^2 cresce per $0 < x < \frac{1}{2}(1 - \sqrt{1 - 2/\mu})$, raggiunge in questo punto il massimo che vale μ , decresce fino a raggiungere nel punto $x = 1/2$ il minimo che vale $\frac{\mu^2}{4}(1 - \mu/4)$, per poi ricrescere e comportarsi nell'intervallo $1/2 \leq x < 1$ in modo simmetrico all'intervallo precedente: g raggiunge il massimo che vale μ nel punto $x = \frac{1}{2}(1 + \sqrt{1 - 2/\mu})$, e poi decresce fino a 0 che viene raggiunto per $x = 1$. Si vede dunque, poiché $\mu > 4$, che l'equazione $x = f^2(x)$ ha esattamente due soluzioni.

Procedendo in questo modo arriviamo a tre conclusioni:

1. A_n consiste di 2^n intervalli aperti e disgiunti. Pertanto $I - (A_0 \cup A_1 \cup \dots \cup A_n)$ consiste di 2^{n+1} intervalli chiusi.
2. f^n trasforma monotonicamente ciascuno di questi intervalli in I .
3. f^n è alternativamente crescente e decrescente su questi intervalli. Dunque il grafico di f avrà esattamente 2^n oscillazioni fra 0 e 1.

Segue da ciò che il grafico della bisettrice $y = x$ incontra il grafico di f^n esattamente in 2^n punti. Detto in altre parole, l'equazione $f^{n+1}(x) = x$ ha esattamente 2^n soluzioni. Quindi f^n ha esattamente 2^n punti fissi, e di conseguenza f ammette esattamente 2^n punti periodici di periodo n in I . Si vede quindi che la struttura della dinamica discreta definita dall'iterazione di f per $\mu > 3$ è assai più complicata di quella che ha luogo per $\mu < 3$. Tanto per cominciare, caratterizziamo l'insieme Λ :

Proposizione 4.7 *Sia $\mu > 2 + \sqrt{5}$. Allora Λ è un insieme di Cantor.*

Dimostrazione. Cominciamo col verificare che se $\mu > 2 + \sqrt{5}$ allora $|f'(x)| > 1 \forall x \in I_0 \cup I_1$. Infatti si ha, risolvendo la disequazione di secondo grado

$$A_0 := \{x \in [0, 1] \mid \mu x(1 - x) > 1\} = \left\{x \mid \frac{1}{2}(1 - \sqrt{1 - 4/\mu}) < x < \frac{1}{2}(1 + \sqrt{1 - 4/\mu})\right\} \quad (4.5)$$

e pertanto

$$I_0 = \left\{x \mid 0 \leq x \leq \frac{1}{2}(1 - \sqrt{1 - 4/\mu})\right\}; \quad I_1 = \left\{x \mid \frac{1}{2}(1 + \sqrt{1 - 4/\mu}) \leq x \leq 1\right\} \quad (4.6)$$

Dato che $f'(x) = -2\mu x + \mu$, si vede subito che $f'(x) > 1$ su I_0 se lo è al suo estremo destro. Sostituendo, questo sarà vero se $\mu(1 - \sqrt{1 - 4/\mu}) > 1$ da cui $\mu > 2 + \sqrt{5}$. Medesimo ragionamento per l'intervallo I_1 . Pertanto, dato che $\Lambda \subset I_0 \cup I_1$, per questi valori di μ esiste $\lambda > 1$ tale che $|f'(x)| > \lambda > 1 \forall x \in \Lambda$. Per la regola di derivazione delle funzioni composte (si veda l'esercizio successivo) da ciò segue anche che $|(f^n)'(x)| > \lambda^n$ su Λ . Ora verifichiamo che Λ non contiene alcun intervallo. Infatti se lo contenesse esisterebbe una coppia di punti $x < y$ in Λ tale che l'intervallo chiuso $[x, y]$ è contenuto in Λ . In tal caso (si veda l'esercizio seguente) $|(f^n)'(\alpha)| > \lambda^n \forall \alpha \in [x, y]$. Prendiamo n in modo tale che sia $\lambda^n|x - y| > 1$. Allora per il teorema del valor medio si può concludere che $|(f^n)(x) - (f^n)(y)| \geq \lambda^n|x - y| > 1$. Ciò implica che almeno uno dei punti $(f^n)(x)$ o $(f^n)(y)$ giace fuori da I , in contraddizione con l'ipotesi. Dunque Λ è totalmente disconnesso.

Poiché Λ è intersezione di intervalli chiusi contenuti gli uni dentro gli altri, è esso stesso un insieme chiuso per un teorema generale di teoria degli insiemi che non è necessario richiamare in dettaglio. Dimostriamo ora che è perfetto. Cominciamo con l'osservare che ogni estremo di ogni intervallo A_k appartiene a Λ . Infatti tali punti vengono ad una qualche iterazione trasformati nel punto fisso a 0, e pertanto l'iterazione non li porta fuori da I . Ora se un qualsiasi $p \in \Lambda$ fosse isolato, ogni punto a lui arbitrariamente vicino dovrebbe uscire da I ad un qualche stadio dell'iterazione di f . Questi punto devono appartenere ad un qualche A_k . Ora si possono dare due casi. O esiste una successione di estremi degli A_k che converge a p , oppure tutti i punti di un intorno di p diversi da p medesimo sono portati fuori da I da una qualche iterazione di f . Nel primo caso non c'è più nulla da dimostrare perché gli estremi degli A_k vengono mandati a 0 e pertanto sono in Λ . Nel secondo caso possiamo assumere che f^n manda p a 0 e tutti gli altri punti dell'intorno a $-\infty$. Ma allora f^n ha un massimo a p e quindi $(f^n)'(p) = 0$. Per la regola di derivazione delle funzioni composte dovrà essere $(f^n)'(f^i(p)) = 0$ per un qualche $i < n$. Pertanto $f^i(p) = 1/2$ perché $1/2$ è il solo punto dove $f'(x) = 0$. Ma allora $f^{i+1}(p) \notin I$ (il valore del massimo a $x = 1/2$ è maggiore di 1) e pertanto $f^n(p) \rightarrow -\infty$ per $n \rightarrow \infty$. Ciò contraddice il fatto che $f^n(p) = 0$, e conclude la dimostrazione della proposizione.

Esercizio 4.29

Sia $f : I \rightarrow \mathbb{R}$ derivabile. Dimostrare che per $x \in I$ vale la formula seguente:

$$(f^n)'(x) = f'((f^{n-1})'(x)) \cdot f'((f^{n-2})'(x)) \cdots f'(x)$$

Osservazione 4.33

Il risultato vale anche per $4 < \mu \leq \mu + \sqrt{5}$, ma la dimostrazione è molto più complicata.

Riassumendo l'analisi fin qui svolta, possiamo dire che abbiamo classificato l'andamento "principale" di tutte le orbite per $\mu > 4$. Vale infatti l'alternativa seguente: o un punto iniziale viene portato a $-\infty$ dalla dinamica discreta generata dall'iterazione di f , oppure la sua intera orbita è contenuta nell'insieme cantoriano Λ . La dinamica fuori da Λ è quindi banale: il punto iniziale semplicemente tende a $-\infty$. La dinamica in Λ è invece sottile ed interessante, e la svilupperemo tramite strumenti matematici assai profondi.

Abbiamo dimostrato che $|f'_\mu(x)| > 1$ su $I_0 \cup I_1$, e ciò implica $|f'_\mu(x)| > 1$ su Λ . Questa condizione equivale ad assumere l'iperbolicità su un intero insieme, e non solo su un punto periodico come l'avevamo definita nel §1 di questo capitolo. In altre parole, possiamo introdurre la nozione di *insieme iperbolico*:

Definizione 4.23 *Un insieme $\Gamma \subset \mathbb{R}$ è iperbolico nonché repulsivo (attrattivo) per f se è chiuso, limitato, invariante rispetto a f (cioè $f(\Gamma) \subset \Gamma$) ed esiste $N > 0$ tale che $|(f^n)'(x)| > 1$ ($|(f^n)'(x)| < 1$) $\forall n \geq N$ e $\forall x \in \Gamma$.*

Esempio 4.29

L'insieme di Cantor Λ per la mappa quadratica quando $\mu > 2 + \sqrt{5}$ è un insieme iperbolico repulsivo con $N = 1$.

5 La dinamica simbolica

Riconsideriamo le successioni con elementi 0 e 1. Denotiamo Σ_2 l'insieme di tutte le successioni s_n i cui elementi assumono solo i valori 0 o 1. I punti di Σ_2 , d'ora in poi denotati $\mathbf{s}$, saranno dunque successioni di elementi 0 e 1. In altre parole:

$$\mathbf{s} \in \Sigma_2 \iff \{\mathbf{s} = (s_0, s_1, s_2, \dots) \mid s_j = 0 \text{ o } 1\}$$

Dunque gli elementi di Σ_2 sono le stringhe infinite del tipo $001001111\dots$ che abbiamo già considerato sotto l'aspetto probabilistico nel Capitolo 2. Vogliamo ora introdurre una nozione di distanza fra due punti $\mathbf{s}$ e $\mathbf{t}$ di Σ_2 . Dati due punti $\mathbf{s}$ e $\mathbf{t}$ in Σ_2 , cioè due successioni $\mathbf{s} = (s_0, s_1, s_2, \dots)$; $\mathbf{t} = (t_0, t_1, t_2, \dots)$; $s_j = 0, 1$; $t_k = 0, 1$ poniamo la seguente

Definizione 5.24 *Dati $\mathbf{s}$ e $\mathbf{t}$ in Σ_2 la loro distanza $d(\mathbf{s}, \mathbf{t})$ è definita nel modo seguente:*

$$d(\mathbf{s}, \mathbf{t}) := \sum_{k=0}^{\infty} \frac{|s_k - t_k|}{2^k} \quad (5.7)$$

Esempio 5.30 *(Calcolo di qualche distanza notevole)*

1. Sia $\mathbf{s} = (0, 0, 0, \dots)$, $\mathbf{t} = (1, 1, 1, \dots)$. Allora:

$$d(\mathbf{s}, \mathbf{t}) = \sum_{k=0}^{\infty} \frac{1}{2^k} = \frac{1}{1 - 1/2} = 2$$

2. Siano $\mathbf{s}$ e $\mathbf{t}$ tali che $s_k = t_k : k = 1, \dots, N$. Allora $d(\mathbf{s}, \mathbf{t}) \leq 2^{-N}$. Infatti:

$$\begin{aligned} d(\mathbf{s}, \mathbf{t}) &= \sum_{k=N+1}^{\infty} \frac{|s_k - t_k|}{2^k} \\ &= \frac{1}{2^{N+1}} \sum_{k=0}^{\infty} \frac{|s_k - t_k|}{2^k} \leq \frac{1}{2^{N+1}} \sum_{k=0}^{\infty} \frac{2}{2^k} = \frac{1}{2^N} \end{aligned}$$

Verifichiamo ora che $d(\mathbf{s}, \mathbf{t})$ ha tutte le proprietà che ci si aspettano da una distanza, e cioè la positività (la distanza di ogni punto da un altro qualsiasi deve essere positiva o al più nulla, e quest'ultimo caso deve valere se e solo se i due punti coincidono), la simmetria (la distanza di un punto da un secondo deve essere uguale a quella del secondo dal primo) e la disuguaglianza triangolare (dati tre punti, la distanza fra due qualsiasi di essi non deve superare la somma delle distanze fra loro stessi ed il terzo punto, come avviene per le distanze fra i vertici di un triangolo: la lunghezza di un lato qualsiasi non supera mai la somma delle lunghezze degli altri due). Per di più, Σ_2 è un insieme limitato rispetto a d . Ciò significa che ogni punto di Σ_2 ha distanza finita da qualsiasi altro. Vale infatti la seguente

Proposizione 5.8

1. $d(\mathbf{s}, \mathbf{t}) \geq 0 \ \forall (\mathbf{s}, \mathbf{t}) \in \Sigma_2$, $d(\mathbf{s}, \mathbf{t}) = 0$ se e solo se $\mathbf{s} = \mathbf{t}$;
2. $d(\mathbf{s}, \mathbf{t}) = d(\mathbf{t}, \mathbf{s})$

3. Vale la disuguaglianza triangolare:

$$d(s, t) \leq d(s, u) + d(u, t)$$

4. Inoltre a distanza massima fra due punti qualsiasi s e t di Σ_2 vale 2.

Dimostrazione Poiché $|s_k - t_k| \leq 1 \forall k$, per la definizione (5.7) di distanza si ha subito

$$d(s, t) \leq \sum_{k=0}^{\infty} \frac{1}{2^k} = 2, \quad \forall (s, t) \in \Sigma_2$$

D'altra parte, è ovvio che $d(s, t) \geq 0 \forall (s, t) \in \Sigma_2$, e poiché (5.7) è una serie a termini non negativi si avrà $d(s, t) = 0$ se e solo se $|s_k - t_k| = 0 \forall k$ cioè se e solo se $s = t$. Infine la disuguaglianza triangolare segue dalla disuguaglianza triangolare numerica: dato che

$$|s_k - t_k| = |s_k - u_k + u_k - t_k| \leq |s_k - u_k| + |u_k - t_k|$$

si ha

$$d(s, t) := \sum_{k=0}^{\infty} \frac{|s_k - t_k|}{2^k} \leq \sum_{k=0}^{\infty} \frac{|s_k - u_k|}{2^k} + \sum_{k=0}^{\infty} \frac{|u_k - t_k|}{2^k} = d(s, u) + d(u, t)$$

e ciò conclude la prova della Proposizione.

Osservazione 5.34

1. Poiché la distanza massima fra due punti qualsiasi di Σ_2 vale 2, Σ_2 è un insieme limitato rispetto alla distanza $d(s, t)$.
2. Sulla retta reale definiamo intervallo aperto di centro a e lunghezza $2b$ l'insieme $I_a := \{x \in \mathbb{R} \mid |x - a| < b\}$, $b > 0$. Nel piano cartesiano definiamo analogamente disco aperto di centro $a = (a_1, a_2)$ e raggio r l'insieme $D_A := \{x = (x_1, x_2) \in \mathbb{R}^2 \mid \sqrt{(x_1 - a_1)^2 + (x_2 - a_2)^2} < r\}$; nello spazio $\mathbb{R}^3$ la palla aperta di centro a è l'insieme

$$P_A := \{x = (x_1, x_2, x_3) \in \mathbb{R}^3 \mid \sqrt{(x_1 - a_1)^2 + (x_2 - a_2)^2 + (x_3 - a_3)^2} < R\}.$$

È dunque naturale definire palla *aperta* di centro a in Σ_2 e raggio R l'insieme:

$$\Pi_a = \{x \in \Sigma_2 \mid d(x, a) < R\}$$

per un qualche $R > 0$.

3. Abbiamo già visto che se i punti $\mathbf{s}$ e $\mathbf{t}$ in Σ_2 hanno le prime N componenti in comune allora si trovano a distanza non superiore a 2^{-N} . Proviamo che anche il viceversa è vero. Infatti, sia

$$d(\mathbf{s}, \mathbf{t}) \leq 2^{-N}$$

D'altra parte si ha

$$d(\mathbf{s}, \mathbf{t}) = \sum_{k=0}^{\infty} \frac{|s_k - t_k|}{2^k} = \sum_{k=0}^N \frac{|s_k - t_k|}{2^k} + \sum_{k=N+1}^{\infty} \frac{|s_k - t_k|}{2^k}$$

Poiché $d(\mathbf{s}, \mathbf{t}) \leq 2^{-N}$, dall'uguaglianza precedente segue subito

$$\frac{|s_k - t_k|}{2^k} \leq 2^{-N}$$

Moltiplicando ambo i membri di quest'ultima per 2^{-N} otteniamo

$$\sum_{k=0}^N 2^{N-k} |s_k - t_k| \leq 1$$

Quest'ultima disuguaglianza può sussistere solo se $s_k = t_k$, $k = 0, \dots, N$.

Questa caratterizzazione ci dice che due successioni in Σ_2 sono tanto più vicine quanto più sono i loro primi termini coincidenti.

Introduciamo ora la nozione alla base di dinamica simbolica, che è una particolare applicazione da Σ_2 in sè (la traslazione, o spostamento, di un'unità⁴) che si itera in modo naturale.

Definizione 5.25 *L'applicazione $\sigma : \Sigma_2 \rightarrow \Sigma_2$ definita da:*

$$\sigma(\mathbf{s}) = \sigma(s_0, s_1, s_2, \dots) := (s_1, s_2, s_3, \dots)$$

si dice traslazione o spostamento unitario.

Osservazione 5.35

1. L'azione della funzione σ su $\mathbf{s} = (s_0, s_1, s_2, \dots)$ è semplicemente la traslazione di un'unità verso sinistra di ciascun elemento della successione, ignorando il primo, cioè s_0 . Ad esempio, se $\mathbf{s} = (0, 1, 1, 0, 1, 0, \dots)$, $\sigma(\mathbf{s}) = (1, 1, 0, 1, 0, \dots)$.

⁴ *Unit shift* in inglese.

2. L'azione della mappa $\sigma \circ \sigma := \sigma^2$ sarà lo spostamento di due unità verso sinistra, e chiaramente per ogni $k = 1, 2, 3, \dots$ si ha

$$\sigma^k(s_0, s_1, s_2, \dots) = (s_k, s_{k+1}, s_{k+2}, \dots); \quad \sigma^0 = Id.$$

Ad esempio, prendendo ancora $\mathbf{s} = (0, 1, 1, 0, 1, 0, \dots)$, e $k = 4$, si ha $\sigma^4 \mathbf{s} = (1, 0, \dots)$.

Data una funzione $f : I \subset \mathbb{R} \rightarrow \mathbb{R}$, sappiamo che essa si dice *continua* in un punto x interno a I se $|f(x) - f(y)| \rightarrow 0$ quando $|x - y| \rightarrow 0$, $(x - y) \in I$. Equivalentemente:

$$\lim_{x \rightarrow y} f(x) = f(y)$$

Una funzione continua in tutti i punti di un intervallo I si dice continua su I . In parole, possiamo dire che una funzione è continua su un intervallo quando la distanza fra i valori assunti in due punti differenti tende a zero al tendere a zero della distanza fra i due punti medesimi. Lo stesso concetto quindi può essere formulato in via del tutto generale una volta che si siano introdotte la nozione di funzione come applicazione fra insiemi ed una nozione di distanza. In altre parole, possiamo porre la seguente

Definizione 5.26 *Data una funzione qualsiasi $f : \Sigma_2 \rightarrow \Sigma_2$ diremo che essa è continua se, per ogni coppia di successioni $(\mathbf{s}, \mathbf{t}) \in \Sigma_2$ tali che $d(\mathbf{s}, \mathbf{t}) < \delta$, esiste $\epsilon(\delta) \rightarrow 0$ per $\delta \rightarrow 0$ tale che*

$$d(f(\mathbf{s}), f(\mathbf{t})) < \epsilon.$$

Osservazione 5.36

Se $f : \Sigma_2 \rightarrow \Sigma_2$, l'immagine $f(\mathbf{s})$ dell'elemento $\mathbf{s}$ sarà la successione $f(\mathbf{s}) := f_0(\mathbf{s}), f_1(\mathbf{s}), f_2(\mathbf{s}), \dots$ con $f_k(\mathbf{s}) = 0, 1 \forall k$. Allora la condizione $d(f(\mathbf{s}), f(\mathbf{t})) < \epsilon$ si scrive in forma esplicita così:

$$d(f(\mathbf{s}), f(\mathbf{t})) = \sum_{k=0}^{\infty} \frac{|f_k(\mathbf{s}) - f_k(\mathbf{t})|}{2^k} < \epsilon \quad (5.8)$$

Possiamo quindi formulare la seguente

Proposizione 5.9 *La funzione $\sigma : \Sigma_2 \rightarrow \Sigma_2$ è continua.*

Dimostrazione Dato $\epsilon > 0$ scegliamo N tale che $2^{-N} < \epsilon$, e sia $\delta := 2^{-N-1}$. Se $\mathbf{s} = (s_0, s_1, s_2, \dots)$, $\mathbf{t} = (t_0, t_1, t_2, \dots)$ soddisfano la condizione $d(\mathbf{s}, \mathbf{t}) < \delta$, per quanto visto sopra si ha $s_k = t_k$ per $0 \leq k \leq N+1$. Dunque i primi 2^N elementi di $\sigma(\mathbf{s})$ e $\sigma(\mathbf{t})$ coincidono. Da ciò segue che $d(\sigma(\mathbf{s}), \sigma(\mathbf{t})) < 2^{-N} = \epsilon$ e ciò conclude la prova della Proposizione.

La dinamica discreta generata dall'iterazione di σ ha molti aspetti interessanti. Vediamone alcuni. Anzitutto denotiamo:

1. $\text{Per}_N(\sigma)$ l'insieme dei punti periodici di periodo N per σ , cioè, ricordiamo l'insieme di tutte le successioni $\mathbf{s} = (s_0, s_1, s_2, \dots)$ tali che $s_{k+N} = s_k \forall k$;
2. $\text{Per}(\sigma) := \bigcup_{N=1}^{\infty} \text{Per}_N(\sigma)$ l'insieme di *tutti* i punti periodici di σ , quindi inclusi i punti fissi.

Allora si ha

Proposizione 5.10

1. $\text{Per}_N(\sigma)$ contiene esattamente 2^N punti (cioè ci sono esattamente 2^N successioni di elementi 0, 1 periodiche di periodo N).
2. L'insieme di tutti i punti periodici $\text{Per}(\sigma)$ è denso in Σ_2 ;
3. σ ammette un'orbita densa in Σ_2 , cioè esiste almeno un punto $\mathbf{r} \in \Sigma_2$ tale che l'insieme formato dai punti $\{\sigma^k \mathbf{r} : k = 0, 1, 2, \dots\}$ è denso in Σ_2 .

Osservazione 5.37

Ricordiamo che, dato un insieme $\mathcal{A}$ munito di una nozione di distanza fra i suoi punti, un suo sottoinsieme $\mathcal{B}$ si dice denso in $\mathcal{A}$ se dato un qualsiasi punto di $\mathcal{A}$ esiste un punto di $\mathcal{B}$ a distanza arbitrariamente piccola da esso. Ad esempio, l'insieme $\mathbb{Q}$ dei numeri razionali è denso nell'insieme $\mathbb{R}$ dei numeri reali perché dato un numero reale qualsiasi esiste sempre una successione di numeri razionali che converge ad esso. Nel caso in esame, l'insieme è Σ_2 , e la distanza è definita dalla (5.7). Quindi la densità di $\text{Per}(\sigma)$ significa che $\forall \mathbf{s} \in \Sigma_2$, e $\forall \epsilon > 0$, esiste $\mathbf{t} \in \text{Per}(\sigma)$ tale che $d(\mathbf{s}, \mathbf{t}) < \epsilon$; allo stesso modo, dire che l'orbita è densa significa dire che dato un punto qualsiasi di Σ_2 esiste almeno un punto dell'orbita a distanza

arbitrariamente piccola da esso; in formule, $\forall \mathbf{s} \in \Sigma_2$, e $\forall \epsilon > 0$, esiste almeno un $p \in \mathbb{N}$ tale che $d(\mathbf{s}, \sigma^p \mathbf{r}) < \epsilon$.

Dimostrazione. I punti periodici di periodo N sono successioni che devono ripetersi nel modo seguente: $\mathbf{s} = (s_0, s_1, \dots, s_{N-1}; s_0, s_1, \dots, s_{N-1}; \dots)$ e quindi ce ne sono 2^N in corrispondenza alle 2^N differenti stringhe di lunghezza N composte dagli elementi 0, 1. Ciò prova il punto 1. Per dimostrare la seconda affermazione, dato un punto qualsiasi $\mathbf{s} = (s_0, s_1, s_2, \dots) \in \Sigma_2$ dobbiamo costruire una successione di punti appartenenti a $\text{Per}(\sigma)$ che vi converga. Definiamo a tale scopo la successione $\mathbf{t}^{(n)} \in \text{Per}(\sigma)$ nel modo seguente:

$$\mathbf{t}^{(n)} := (s_0, \dots, s_{n-1}; s_0, \dots, s_{n-1}; \dots)$$

e facciamo vedere che $\lim_{n \rightarrow \infty} \mathbf{t}^{(n)} = \mathbf{s}$. In altre parole, per l'espressione esplicita (5.8) dobbiamo far vedere quanto segue:

$$\lim_{n \rightarrow \infty} \sum_{k=0}^{\infty} \frac{|\mathbf{t}_k^{(n)} - s_k|}{2^k} = 0.$$

$\forall n$ costruiamo una successione i cui primi n elementi coincidono i primi n elementi della successione data e poi la prolunghiamo periodicamente. Poiché i primi n elementi di $\mathbf{s}$ e $\mathbf{t}_n$ coincidono sappiamo che $d(\mathbf{s}, \mathbf{t}_n) < 2^{-n}$; facendo $n \rightarrow \infty$ otteniamo il risultato. Per provare l'ultima affermazione dobbiamo ancora approssimare $\mathbf{s}$ con una successione, stavolta però appartenente ad un orbita. A tale scopo, consideriamo il seguente dato iniziale

$$\mathbf{u} = (01 \mid 00 \ 01 \ 10 \ 11 \mid 000 \ 001 \ 010 \ \dots)$$

che viene costruito mettendo in fila l'uno dopo l'altro prima tutte le possibili stringhe di 1 elemento, poi di 2 elementi, poi di 3 elementi, ecc. È chiaro che un qualche iterato di $\mathbf{u}$ genererà una successione che coincide con ogni successione data in un numero arbitrariamente grande di elementi, da cui la densità. Ciò conclude la prova della proposizione.

6 Mappa logistica sull'insieme di Cantor e dinamica simbolica

In questo paragrafo faremo vedere che l'azione della mappa f sull'insieme di Cantor Λ è identica all'azione di σ su Σ_2 . Per far ciò dobbiamo introdurre la nozione di dinamica

simbolica, e poi sfruttare una nozione matematica molto generale, l'omeomorfismo, di cui ricordiamo la definizione.

Definizione 6.27 *Dati due insiemi A e B sui quali sono definite operazioni di composizione interna, diremo che l'applicazione h definisce un omeomorfismo h fra A e B se:*

1. h è iniettiva;
2. h conserva le operazioni, cioè se l'immagine attraverso h della composizione di due elementi in A è l'operazione di composizione in B che agisce sugli elementi immagine.

Esempio 6.31

L'esponenziale complesso $x \mapsto e^x$ definisce un omeomorfismo fra la retta reale $\mathbb{R}$ e la semiretta reale positiva $\mathbb{R}_+$ che trasforma l'addizione in moltiplicazione: l'applicazione $x \mapsto e^{-x}$ è infatti iniettiva, e per di più il corrispondente dell'elemento $x + y$ è l'elemento $e^{x+y} = e^x \cdot e^y$.

Per applicare questo concetto alla mappa logistica, cominciamo col ricordare che $\Lambda \subset I_0 \cup I_1$. Se $x \in \Lambda$, tutti i punti della sua orbita appartengono a Λ e quindi appartengono a I_0 oppure a I_1 dato che questi due intervalli sono disgiunti.

Definizione 6.28 *Sia f_μ la mappa logistica. Si dice itinerario del punto iniziale x la successione $S(x) = (s_0, s_1, s_2, \dots)$ definita nel modo seguente*

$$s_k = \begin{cases} 0 & \text{se } f_\mu^k \in I_0 \\ 1 & \text{se } f_\mu^k \in I_1 \end{cases} \quad (6.9)$$

Osservazione 6.38

L'azione di S stabilisce un'applicazione fra Λ e Σ_2 in più di un modo analoga alla decomposizione binaria di un numero reale. Infatti, tanto per cominciare, ad ogni traiettoria, che ha sempre luogo in Λ se il punto iniziale vi appartiene, viene associata una successione di 0 e 1, cioè di due simboli. Le domande da porsi sono:

1. L'applicazione S sarà una codificazione? In altre parole, ci sarà una corrispondenza biunivoca fra gli elementi di Λ e quelli di Σ_2 (simile a quella fra numeri frazionari e loro sviluppi binari)?

2. L'applicazione S sarà continua, cioè manderà elementi "vicini" in Λ in elementi vicini in Σ_2 ?
3. Sarà possibile codificare anche in Σ_2 l'azione della dinamica discreta generata dall'iterazione di f_μ su Λ ? In altre parole, quale sarà l'immagine in Σ_2 della dinamica discreta su Λ ? Essa sarà una dinamica nello spazio dei simboli, e pertanto verrà detta *dinamica simbolica*. L'idea è che la dinamica simbolica sia assai più facile da descrivere di quella originale, e quindi permetta di capire quest'ultima se la "codificazione" nello spazio Σ_2 è fedele.

A queste domande dà risposta il seguente

Teorema 6.7 *Se $\mu > 2 + \sqrt{5}$ allora l'applicazione $S : \Lambda \rightarrow \Sigma_2$ è un omeomorfismo topologico che coniuga la dinamica discreta definita dall'iterazione di f_μ su Λ con la dinamica definita dall'iterazione di σ su Σ_2 .*

Osservazione 6.39

1. In questo caso dire che S è un omeomorfismo topologico significa semplicemente che S gode delle seguenti proprietà:
 1. S è un'applicazione iniettiva;
 2. S è continua assieme a S^{-1} . La continuità di S significa che se $(x, y) \in \Lambda$ sono tali che $|x - y| < \epsilon$ allora esiste $\delta(\epsilon) \rightarrow 0$ per $\epsilon \rightarrow 0$ tale che $d(S(x), S(y)) < \delta(\epsilon)$, dove d è definita dalla (5.7). La continuità di S^{-1} significa che se $(s, t) \in \Sigma_2$ sono tali che $d(s, t) < \epsilon$, allora esiste $\delta(\epsilon) \rightarrow 0$ per $\epsilon \rightarrow 0$ tale che $|x - y| < \delta(\epsilon)$ dove $x = S^{-1}(s)$ e $y = S^{-1}(t)$.
2. La dinamica discreta definita dall'iterazione di f su Λ e quella definita dall'iterazione di σ su Σ_2 si dicono coniugate (o topologicamente coniugate) se

$$S \circ f_\mu = \sigma \circ S \quad (6.10)$$

cioè se l'azione della dinamica commuta con l'omeomorfismo. In altre parole, scrivendo la (6.10) più esplicitamente, le due dinamiche sono coniugate se

$$S(f_\mu(x)) = \sigma S(x) \quad (6.11)$$

da cui $S(f_\mu^k(x)) = \sigma^k S(x)$, $k = 1, 2, \dots$. In altre parole: le due dinamiche sono coniugate dall'omeomorfismo S se le orbite di f_μ in A di punto iniziale x sono mandate nelle orbite di σ in Σ_2 di punto iniziale $S(x)$.

3. Due modi equivalenti di riscrivere la (6.10) sono

$$f_\mu = S^{-1} \circ \sigma \circ S; \quad S \circ f_\mu \circ S^{-1} = \sigma \quad (6.12)$$

e cioè, esplicitamente:

$$f_\mu(x) = S^{-1}(\sigma S(x)), \quad \sigma S(x) = S(f_\mu(S^{-1}(x)))$$

Dimostrazione.

È riportata in Appendice per completezza, essendo piuttosto complicata.

Se ad esempio riconsideriamo l'omeomorfismo fra $(\mathbb{R}; +)$ e $(\mathbb{R}_+, \times)$ definito dall'applicazione esponenziale $h := x \mapsto e^x$, e definiamo su $\mathbb{R}$ la traslazione $T_\alpha := x + \alpha$, l'operazione su $\mathbb{R}_+$ topologicamente coniugata a T_α è la moltiplicazione $R_\alpha := e^x \mapsto e^{\alpha+x}$.

In generale si pone la definizione seguente:

Definizione 6.29 *Siano $f : A \rightarrow A$ e $g : B \rightarrow B$ due applicazioni. Diremo che f e g sono (topologicamente) coniugate se esiste un omeomorfismo topologico $h : A \leftrightarrow B$ tale che $h \circ f = g \circ h$, cioè $h(f(x)) = g(h(x))$.*

Esercizio 6.30

Dimostrare che la coniugazione topologica gode delle proprietà di

1. Identità: f è topologicamente coniugata con sè stessa;
2. Riflessività: se f è topologicamente coniugata a g , g è topologicamente coniugata a f ;
2. Transitività: se f è topologicamente coniugata a g , e g è topologicamente coniugata a h , f è topologicamente coniugata a h .

Esempio 6.32

Sia $Q_c := x^2 + c$, $x \in \mathbb{R}$. Facciamo vedere che, se $c < \frac{1}{4}$, Q_c è (topologicamente) coniugata alla mappa logistica $\mu x(1-x)$ per un ben determinato $\mu > 1$ tramite l'omeomorfismo lineare di $\mathbb{R}$ $h(x) = \alpha x + \beta$.

Infatti la condizione $Q_c(h(x)) = h(\mu x(1-x))$ si scrive esplicitamente così:

$$(\alpha x + \beta)^2 + c = \alpha \mu x(1 - \mu x) + \beta \implies \alpha^2 x^2 + 2\alpha\beta x = \beta^2 + c = \alpha \mu x - \alpha \mu^2 x^2 + \beta$$

da cui, per il principio di identità dei polinomi:

$$\alpha^2 = -\alpha \mu^2; \quad 2\alpha\beta = \alpha \mu; \quad \beta^2 + c = \beta$$

Poiché μ deve essere reale dovremo scegliere $\alpha < 0$. Allora si trova $\mu = |\alpha|$, $\beta = |\alpha|/2$, e $|\alpha|$ medesimo è determinato dalla terza condizione, che genera l'equazione di secondo grado $|\alpha|^2 - 2|\alpha| + 4c = 0$. Questa equazione deve avere radici reali affinché h sia un omeomorfismo di $\mathbb{R}$. La condizione affinché le abbia è appunto $1 - 4c > 0$ da cui $c < \frac{1}{4}$.

Osservazione 6.40

Applicazioni (topologicamente) coniugate generano tramite iterazione dinamiche discrete completamente equivalenti. Ad esempio:

1. Se f è (topologicamente) coniugata a g tramite l'omeomorfismo h , e p è un punto fisso di f , allora $h(p)$ è un punto fisso di g . Infatti se $f(p) = p$, allora $g(h(p)) = h(p)$ perché $h(f(x)) = g(h(x))$.
2. Allo stesso modo h mette in corrispondenza biunivoca $\text{Per}_n(f)$ e $\text{Per}_n(g)$, e quindi anche $\text{Per}(f)$ e $\text{Per}(g)$.
3. Punti successivamente periodici e orbite asintotiche di f vengono trasformati da h in punti successivamente periodici e orbite asintotiche di g .

In particolare, dato che la mappa logistica f_μ è (topologicamente) coniugata alla traslazione di un posto σ , possiamo immediatamente concludere che f_μ gode delle sorprendenti proprietà dimostrate così agevolmente per σ . In altre parole, vale il seguente

Teorema 6.8 *Sia $f_\mu(x) = \mu x(1-x)$, $\mu > 2 + \sqrt{5}$, $x \in [0, 1]$. Allora*

1. *L'insieme $\text{Per}_n(f_\mu)$ contiene esattamente 2^n punti;*
2. *L'insieme $\text{Per}(f_\mu)$ è denso in Λ ;*

3. f_μ ammette un'orbita densa in Λ .

Osservazione 6.41

1. Le medesime affermazioni del teorema valgono anche per la mappa Q_c su $\mathbb{R}$, $c < 1/4$, perché essa è topologicamente coniugata alla mappa logistica tramite l'omeomorfismo lineare $h(x)$.
2. Questo teorema mette in chiaro la potenza della dinamica simbolica; più in generale, fa vedere come un problema apparentemente impossibile da affrontare diventa quasi banale se affrontato nelle coordinate "giuste". Naturalmente la difficoltà viene spostata nella ricerca di queste coordinate, cioè nella costruzione dell'omeomorfismo h .

7 Dipendenza delicata dalle condizioni iniziali. Caos

La mappa quadratica mette in luce in maniera stupefacentemente chiara un fenomeno molto importante e a tutt'oggi dimostrato rigorosamente solo in una piccola frazione dei casi nei quali si manifesta numericamente⁵, cioè la dipendenza estremamente delicata della dinamica dalle condizioni iniziali. Ci sono molte definizioni possibili di comportamento caotico, o caos, nell'evoluzione dinamica, la cui origine sta appunto nella dipendenza delicata dalle condizioni iniziali. Qui, non potendo ancora fare uso di concetti di teoria della misura o di teoria della probabilità come in teoria ergodica, si seguirà un approccio "topologico". Cominciamo con la seguente definizione:

Definizione 7.30 *La funzione $f : I \rightarrow J$ si dice (topologicamente) transitiva se per ogni coppia di insiemi aperti $U \subset J$, $V \subset J$ esiste $k > 0$ tale che $f^k(U) \cap V \neq \emptyset$.*

Intuitivamente, una mappa topologicamente transitiva è fatta in modo tale che il fascio di orbite che escono da un intorno arbitrariamente piccolo prima o poi visita *qualsiasi* altro intorno, cioè ha intersezione non vuota con quest'ultimo. Si noti che:

⁵La comprensione di questo fenomeno non è ancora completa anche se i concetti fondamentali che definiscono il comportamento caotico e l'isolamento del meccanismo che lo determina risalgono a più di un secolo fa, e precisamente all'opera di Henri Poincaré (Nancy 1854-Parigi 1912), Professore di Fisica matematica alla Sorbona a Parigi. Poincaré fu anche un grandissimo epistemologo. È vivamente consigliata la lettura dei suoi libri *La scienza e l'ipotesi*, *Il valore della scienza* e *La scienza e il metodo*, scritti in uno stile quanto mai accessibile.

1. Se una mappa ammette due insiemi invarianti disgiunti *non può essere* topologicamente transitiva, perché le orbite che hanno origine nel primo insieme, dovendo rimanervi, non potranno mai raggiungere il secondo e viceversa.
2. Ugualmente, se la mappa ammette solo orbite periodiche dello stesso periodo *non può essere* topologicamente transitiva, perché si può dimostrare che nessun fascio diverso dall'insieme stesso potrà visitare l'intero insieme se il periodo è il medesimo per ogni condizione iniziale;
3. Se invece la mappa ammette anche una sola orbita densa chiaramente è topologicamente transitiva. Infatti la densità implica che ogni aperto contiene almeno un punto dell'orbita, e quindi, ancora per la densità, il fascio di orbite che escono da ogni aperto prima o poi visiterà qualsiasi altro aperto.

Il concetto successivo è quello di *dipendenza delicata*⁶ dalle condizioni iniziali.

Definizione 7.31

1. Si dice che la dinamica generata dall'iterazione di $f : J \rightarrow I$ ha dipendenza delicata dalle condizioni iniziali quando esiste $\delta > 0$ tale che, per ogni $x \in J$ e ogni intorno $\mathcal{N}$ di x , esistono $y \in \mathcal{N}$ e $n > 0$ tali che $|f^n(x) - f^n(y)| > \delta$.
2. Si dice che la dinamica generata dall'iterazione di $f : J \rightarrow I$ è espansiva se esiste $\nu > 0$ tale che, per ogni $(x, y) \in J, x \neq y$, esiste n tale che $|f^n(x) - f^n(y)| > \nu$.

Osservazione 7.42

1. Si noti che la nozione di espansività è strettamente più forte di quella di dipendenza delicata dalle condizioni iniziali. Infatti in questo caso, fissato un dato iniziale, *tutti* i punti iniziali arbitrariamente vicini se ne allontaneranno al trascorrere del tempo discreto n . Affinché si verifichi la dipendenza delicata invece è sufficiente che esista anche un solo dato iniziale in ogni intorno di x che se ne allontana al trascorrere del tempo.

⁶Mi permetto di proporre questa traduzione dell'espressione inglese *sensitive dependence on initial conditions*, perché la traduzione consueta, dipendenza sensibile, mi sembra un esempio dei disgustosi anglicismi che ogni giorno di più funestano l'italiano.

2. La dipendenza delicata dalle condizioni iniziali esclude nel modo più assoluto la stabilità. Infatti, scelto a piacere il punto iniziale x , esistono punti arbitrariamente vicino a lui che prima o poi se ne allontanano di almeno δ , e quindi sicuramente la condizione iniziale x non è stabile. Si noti ancora una volta che *non tutti* i punti vicini ad x devono separarsi da lui in seguito alla dinamica discreta, ma ce ne deve essere almeno uno *in ogni* intorno di x .

Un'osservazione molto importante sul piano concettuale, e forse ancor più sul piano concreto, è la seguente:

se una mappa possiede la proprietà della dipendenza delicata dalle condizioni iniziali la dinamica da essa generata non è in pratica calcolabile numericamente.

Infatti gli arrotondamenti introdotti necessariamente dal calcolatore vengono amplificati in modo arbitrariamente grande dall'iterazione. In altre parole, se lavoriamo ad esempio con 16 cifre esatte il calcolatore assumerà uguali due numeri che differiscono solo per la 17-sima cifra decimale. Quindi se eseguiamo l'iterazione di ambedue tramite il calcolo numerico otterremo la medesima traiettoria; tuttavia può essere benissimo, se la dinamica è caotica, che le loro traiettorie *vere* siano totalmente differenti. La stessa cosa succederà se invece di 16 cifre esatte lavoriamo con 32, ecc. Non possiamo sfuggire da questo problema perché la precisione dei calcoli impostabili su *qualsiasi* calcolatore rimane finita. La conclusione è la seguente: pur essendo il sistema completamente deterministico, nel senso che l'evoluzione dinamica di ogni assegnato dato iniziale è in linea di principio univocamente determinata, la dinamica stessa può risultare così complicata nella sua dipendenza dalle condizioni iniziali che le sole previsioni possibili per l'andamento a tempi lunghi (cioè per $n \rightarrow \infty$) sono quelle basate sui metodi probabilistici. A questo fenomeno si dà il nome di *impredicibilità*.

(Osservazione importante per chi studierà in futuro la meccanica quantistica: questo fenomeno non va assolutamente confuso con l'indeterminismo intrinseco nella meccanica quantistica).

Esempio 7.33

La mappa logistica $f_\mu(x) = \mu x(1 - x)$ con $\mu > 2 + \sqrt{5}$ dipende in maniera delicata dalle condizioni iniziali su Λ . Infatti si scelga δ minore della lunghezza $\sqrt{1 - \frac{4}{\mu}}$ di A_0 , dove A_0 è l'apertura fra gli intervalli I_0 e I_1 (si ricordino le formule (4.5,4.6) nella dimostrazione del Teorema 4.7). Siano $(x, y) \in \Lambda$ con $x \neq y$. Allora $S(x) \neq S(y)$, e quindi gli itinerari di x e y devono essere differenti in almeno un posto, ad esempio l' n -esimo. Se così è, però, $f_\mu^n(x)$ e $f_\mu^n(y)$ devono trovarsi da parti opposte rispetto ad A_0 , cosicché

$$|f_\mu^n(x) - f_\mu^n(y)| > \frac{1}{3} = l(A_0) > \delta$$

(in altre parole: all'istante n la dinamica ha fatto evolvere le due condizioni iniziali in modo da trovarsi da parti opposte dell'intervallo centrale) e questo prova la dipendenza delicata dalle condizioni iniziali.

Esempio 7.34

La rotazione della circonferenza di angolo irrazionale della Definizione (4.17) genera una dinamica discreta topologicamente transitiva, dato che ogni orbita è densa, ma la sua dipendenza dalle condizioni iniziali non è delicata: trattandosi di una traslazione, due condizioni iniziali separate da una certa distanza angolare rimangono alla medesima distanza ad ogni iterazione:

$$T_\lambda^n(\theta_1) - T_\lambda^n(\theta_2) = \theta_1 + 2\pi n\lambda - \theta_2 - 2\pi n\lambda = \theta_1 - \theta_2.$$

Siamo ora in grado di dare la definizione di dinamica discreta caotica generata dall'iterazione di una mappa o, per abbreviare, la definizione di mappa caotica.

Definizione 7.32 *Sia $V \subset \mathbb{R}$ un intervallo chiuso. La mappa $f : V \rightarrow V$ si dice caotica se gode delle seguenti tre proprietà:*

1. *f ha una dipendenza delicata dalle condizioni iniziali;*
2. *f è topologicamente transitiva;*
3. *I punti periodici di f sono densi in V .*

Osservazione 7.43

1. Questa definizione continua a valere senza alcuna variazione nel caso in cui V sia un insieme arbitrario nel quale sia definita una distanza, e f una mappa da V in sè.
2. Secondo questa definizione gli ingredienti per fabbricare una mappa caotica sono tre: l'impredicibilità, l'indecomponibilità, ed un elemento di regolarità. L'impredicibilità è una conseguenza diretta, come abbiamo visto, della dipendenza delicata dalle condizioni iniziali. La transitività topologica impedisce alla dinamica di essere decomposta in sottoinsiemi invarianti più semplici dai quali non esce. In altre parole, un sistema caotico è *indecomponibile*. Infine esiste anche un elemento di regolarità in mezzo a tutto questo comportamento casuale, e cioè la presenza di un insieme denso di punti periodici.
3. (*Per chi riprenderà in mano queste note dopo avere proseguito negli studi di matematica*). Questi tre ingredienti sono definibili per sistemi ben più generali e importanti delle mappe, ad esempio le equazioni differenziali generate dalle leggi della meccanica. Le mappe unidimensionali non lineari sono semplicemente gli esempi più semplici dove queste definizioni possono essere riscontrate, e servono come utile modello per il comportamento caotico dei sistemi fisici veri che sono molto più complicati.

Esempio 7.35

1. La mappa $f : S^1 \rightarrow S^1$ definita da $f(\theta) := 2\theta$, già considerata nell'esempio 3.20, è caotica. Infatti: la distanza angolare fra due punti iniziali si raddoppia ad ogni iterazione. Pertanto f è espansiva ed in particolare gode della proprietà della dipendenza delicata dalle condizioni iniziali. La transitività topologica è anch'essa conseguenza dell'espansività: ogni arco in S^1 di ampiezza arbitrariamente piccola viene prima o poi amplificato da una qualche f^k in modo da coprire tutto S^1 e quindi in particolare ogni altro arco in S^1 . Infine abbiamo visto nell'esempio 4.22, n.4, che l'insieme dei punti periodici della mappa $\theta \mapsto 2\theta$ è denso, e quindi la mappa è caotica.
2. La famiglia delle mappe logistiche $f_\mu(x) = \mu x(1 - x)$ è caotica per $\mu > 2 + \sqrt{5}$.

Il secondo esempio ha una marcata differenza dal primo: il comportamento caotico non avviene per qualsiasi scelta dei dati iniziali, ma solo se questi ultimi vengono scelti in un opportuno e "piccolo" sottoinsieme del dominio $I = [0, 1]$, cioè l'insieme di Cantor A .

Facciamo tuttavia vedere che per il valore particolare $\mu = 4$ la mappa logistica è in effetti caotica su tutto $[0, 1]$.

Esercizio 7.31

Si dimostri che $f_4(x) := 4x(1 - x)$ è caotica sull'intervallo $I = [0, 1]$.

Soluzione. Usiamo ancora lo strumento della coniugazione. Sia ancora $g(\theta) := \theta \mapsto 2\theta$ la mappa di S^1 in sé riconsiderata sopra. Definiamo l'applicazione $h_1 : S^1 \rightarrow [-1, 1]$ nel modo seguente: $h_1 := \cos(\theta)$. In altre parole, h_1 proietta θ sull'asse x . Poiché h_1 è una mappa $2-1$ (salvo per $\theta = \pi/2$) dato che $\cos \theta = \cos(-\theta)$, essa non genera un omeomorfismo. Tuttavia procediamo lo stesso a costruire una coniugazione. Sia $q(x) = 2x^2 - 1$. Allora abbiamo

$$h_1 \circ g(\theta) = \cos(2\theta) = 2\cos^2 \theta - 1 = q \circ h_1(\theta)$$

cosicché h_1 coniuga g con q . La coniugazione è topologica per l'evidente continuità di h_2 . Ora facciamo vedere che q è topologicamente coniugata a f_4 . Sia infatti $h_2(t) := \frac{1-t^2}{2}$. Allora si ha

$$(f_4 \circ h_2)(t) = 2(1-t)\frac{t+1}{2} = 2t^2 - 1; \quad (h_2 \circ q)(x) = -\frac{1}{2} + 2x^2 - \frac{1}{2} = 2x^2 - 1$$

Dunque $f_4 \circ h_2 = h_2 \circ q$. Abbiamo già osservato che la coniugazione topologica gode della proprietà transitiva: se f è coniugata a q tramite h_1 , e q è coniugata a g tramite h_2 , f è coniugata a g tramite la composizione $h_1 \circ h_2$. Pertanto f_4 è topologicamente coniugata a g tramite $h_1 \circ h_2$, mappa continua perché lo sono tanto h_1 quanto h_2 . Poiché g è caotica su tutto S^1 , F_4 lo è su tutto I . Rimane solo da discutere il fatto che h_1 non è un omeomorfismo. Questo punto lo lasciamo al lettore.

Capitolo 5

Biforcazioni e transizione al caos

Abbiamo visto che la mappa logistica si comporta in modo drammaticamente differente al crescere del parametro μ . Da una situazione, per $\mu < 3$, di massima semplicità in cui tutte le orbite convergono verso il punto fisso p_μ si passa se $\mu > 4$ ad una situazione caotica. Quale meccanismo genera il caos al crescere di μ ? Abbiamo visto che a $\mu = 3$ ha luogo un fenomeno di biforcazione del punto fisso. Ebbene faremo vedere che proprio il moltiplicarsi e l'accumularsi di simili fenomeni, in particolare le biforcazioni con raddoppiamento di periodo dei punti periodici, è il meccanismo che genera il caos. Poiché la trattazione matematica completa di questo punto richiede conoscenze più avanzate, ci limiteremo a "provare" numericamente questo fenomeno, dopo avere introdotto però le nozioni matematiche necessarie per descriverlo.

1 Esempi di biforcazioni

Cominciamo al solito con qualche esempio, scegliendone due molto significativi perché rappresentano i due casi fondamentali.

Esempio 1.36 (*Biforcazione tangenziale, o nodo-sella*)

Consideriamo la famiglia di funzioni $E_\lambda(x) = \lambda e^x$, $x > 0$, al variare del parametro $\lambda > 0$. Facciamo vedere che questa famiglia subisce una biforcazione del suo punto fisso quando $\lambda = \frac{1}{e}$. Per capirlo geometricamente basta dare un'occhiata ai grafici della funzione per $\lambda > 1/e$, $\lambda = \frac{1}{e}$ e $\lambda < 1/e$ rispettivamente (figg. 5.1 a,b,c).

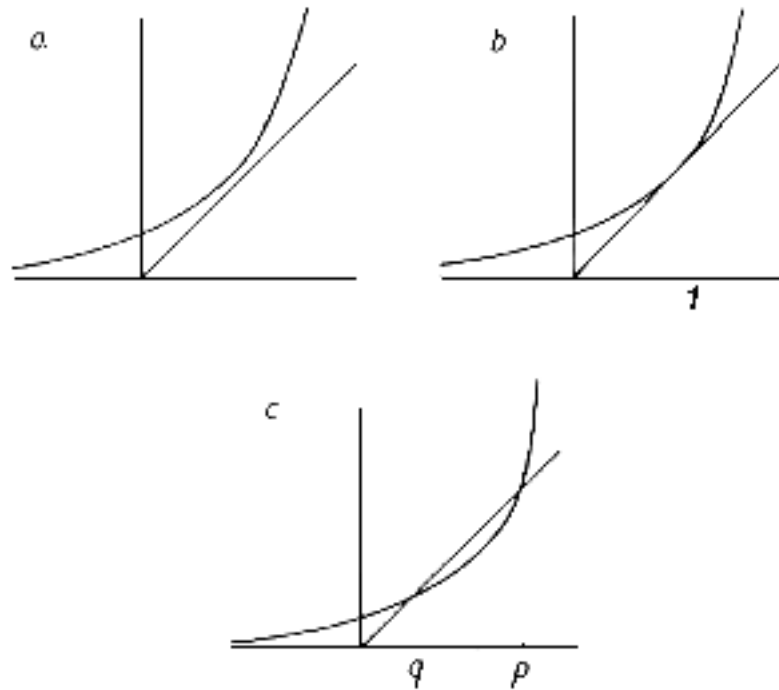


Figure 5.1: I grafici di $E_\lambda(x) = \lambda e^x$ per (a) $\lambda > \frac{1}{e}$, (b) $\lambda = \frac{1}{e}$ e (c) $\lambda < 1/e$

Per $\lambda > \frac{1}{e}$ si ha $e^x > x$, e quindi non ci sono punti fissi. Per $\lambda = \frac{1}{e}$ la retta $y = x$ è tangente alla curva $y = \lambda e^x$ nel punto $x = 1, y = 1$. Per $\lambda < 1/e$ la bisettrice $y = x$ seca la curva $y = \lambda e^x$ in due punti: $q < 1$ con $E'_\lambda(q) < 1$ e a $p > 1$ con $E'_\lambda(p) > 1$. Dunque ci sono due punti fissi per $\lambda < 1/e$: al decrescere del parametro, nascono due punti fissi quando λ passa per $\frac{1}{e}$. Non è difficile tracciare l'intero diagramma di fase. Comunque proviamo le proprietà seguenti della dinamica discreta $E_\lambda^n(x)$ generata da $E_\lambda(x)$:

1. Se $\lambda > \frac{1}{e}$ $E_\lambda^n(x) \rightarrow \infty$ per ogni x quando $n \rightarrow \infty$.
2. Se $\lambda = \frac{1}{e}$, $E_\lambda(1) = 1$. Se $x < 1$, $E_\lambda^n(x) \rightarrow 1$ per $n \rightarrow \infty$. Se $x > 1$, $E_\lambda^n(x) \rightarrow +\infty$ per $n \rightarrow \infty$.
3. Se $0 < \lambda < \frac{1}{e}$, $E_\lambda(q) = q$ e $E_\lambda(p) = p$. Se $x < p$, $E_\lambda^n(x) \rightarrow q$ per $n \rightarrow \infty$; se $x > p$, $E_\lambda^n(x) \rightarrow +\infty$ per $n \rightarrow \infty$.

La proprietà 1 è banale perché se $\lambda > \frac{1}{e}$ si ha chiaramente $E_\lambda^n(x) > 1$ per ogni $x > 0$. Anche la proprietà 2 è evidente: se $\lambda = \frac{1}{e}$ $E_\lambda(1) = 1$ per costruzione; per di più $E_\lambda(x) < 1$ per $x < 1$ da cui segue subito che $E_\lambda^n(x) \rightarrow e^0 = 1$ per $n \rightarrow \infty$. D'altra parte, $E_\lambda(x) > 1$ per $x > 1$ e quindi $E_\lambda^n(x) \rightarrow +\infty$ per $n \rightarrow \infty$. Se $0 < \lambda < \frac{1}{e}$ le proprietà di punto fisso sono vere per costruzione. Se $x < p$, $\lambda e^x < p$ e quindi la successione $E_\lambda^{on}(x)$ è decrescente, strettamente perché $E'_\lambda(x) > 0$. Dunque necessariamente $E_\lambda^{on}(x) \rightarrow q = E_\lambda(q) < 1$ per $n \rightarrow \infty$.

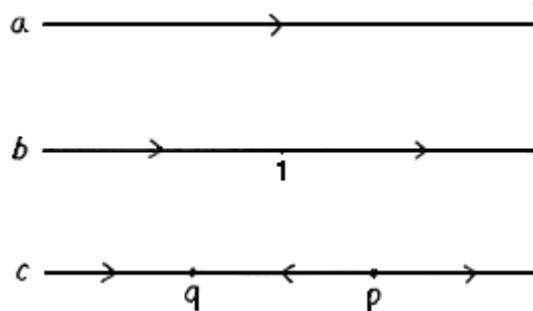


Figure 5.2: Diagramma di fase di λe^x . a: $\lambda > \frac{1}{e}$. b: $\lambda = \frac{1}{e}$. c: $\lambda < \frac{1}{e}$

È utile introdurre anche il *diagramma di biforcazione* cioè la figura seguente:

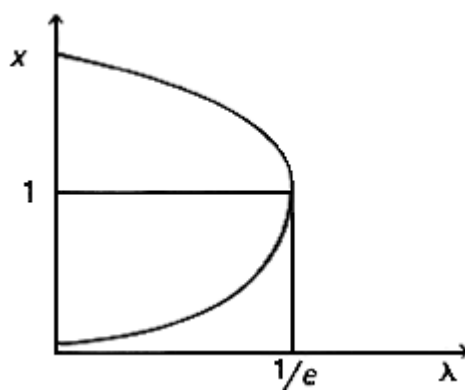


Figure 5.3: Diagramma di biforcazione di λe^x . x come funzione di λ . In ascissa: il parametro. In ordinata: le ascisse dei punti fissi (o periodici) biforcanti. Tagliando il grafico con una retta verticale a $\lambda = \lambda_0$ si ottengono le ascisse dei punti biforcanti per il particolare valore λ_0 del parametro.

Esempio 1.37 (*Biforcazione che raddoppia il periodo*)

Consideriamo ancora la famiglia $E_\lambda(x) = \lambda e^x$, stavolta per $\lambda < 0$. Se $\lambda = -e$, $E_\lambda(-1) = -1$ e $E'_\lambda(-1) = -1$. Dunque -1 è un punto fisso di E_λ non iperbolico. Se $\lambda > -e$ si può dimostrare (vedasi l'esercizio seguente) che il punto fisso di E_λ è attrattivo; quando $\lambda < -e$ è repulsivo. Pertanto la naura della dinamica con condizioni iniziali vicine al punto fisso subisce un marcato cambiamento a $\lambda = -e$. Per di più questo non è tutto. Consideriamo E_λ^2 . Usando il calcolo infinitesimale è facile vedere (esercizio successivo) che E_λ^2 è una funzione crescente, con la concavità verso l'alto se $E_\lambda(x) < -1$ e verso il basso se $E_\lambda(x) > -1$. Pertanto E_λ^2 acquista 2 nuovi punti fissi a q_1 e q_2 quando λ diminuisce sotto $-e$. Poiché sappiamo che $E_\lambda(x)$ ha un solo punto fisso, questi due punti dovrenno essere punti periodici di periodo 2 per $E_\lambda(x)$. Dinamicamente, cioè al variare del parametro, vediamo che la biforcazione che raddoppia il periodo comporta quanto segue:

1. Un punto fisso che da attrattivo diventa repulsivo;
2. La nascita di una nuova orbita di periodo 2.

Nell'esempio precedente, per di più, notiamo che l'attrattività persa dal punto fisso viene ereditata dall'orbita periodica.

Esercizio 1.32

Dimostrare le affermazioni non dimostrate nell'esempio precedente.

Soluzione. Anzitutto notiamo che c'è un solo punto fisso perché l'equazione $\lambda e^x = x$ ha evidentemente una ed una sola soluzione x^* per $\lambda < 0$, $x < 0$. Inoltre, se $\lambda < -e$ si ha $|E_\lambda(x^*)| < 1$ mentre $|E_\lambda(x^*)| > 1$ se $\lambda > -e$. Da qui l'attrattività e la repulsività nei due casi, rispettivamente. Consideriamo poi la funzione $E^2(x) = \lambda \exp(\lambda \exp x)$. Si ha:

$$\begin{aligned}\frac{dE^2(x)}{dx} &= \lambda^2 \exp(\lambda \exp x) > 0 \quad \forall x \in \mathbb{R} \\ \frac{d^2 E^2(x)}{dx^2} &= \lambda^2(1 + \lambda) \exp x \exp(\lambda \exp x)\end{aligned}$$

da cui

$$\frac{d^2 E^2(x)}{dx^2} > 0 \quad \text{per } \lambda < -1; \quad \frac{d^2 E^2(x)}{dx^2} < 0 \quad \text{per } \lambda > -1.$$

Il diagramma di biforcazione si trova in fig.5.4.

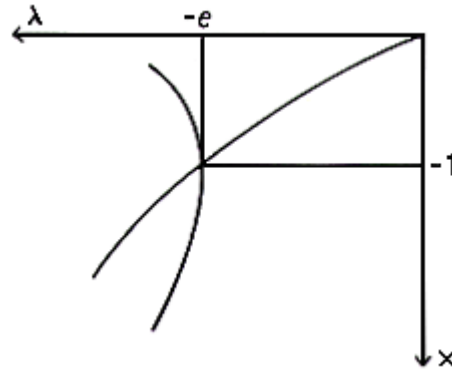


Figure 5.4: Diagramma di biforcazione con raddoppio di periodo per $E_\lambda(x) = \lambda e^x$. x è graficato come funzione di λ .

2 Biforcazioni: nodo-sella e raddoppio di periodo

Enunciamo qui i teoremi generali riguardo le biforcazioni che ci serviranno in seguito. Le dimostrazioni sono relegate in Appendice.

Il primo risultato fa vedere che finché il punto fisso si mantiene iperbolico non potranno mai esserci biforcazioni.

Teorema 2.9 *Sia $f_\lambda(x)$ una famiglia ad un parametro di funzioni regolari da un intervallo aperto $I \subset \mathbb{R}$ in $\mathbb{R}$; si assuma poi che il parametro λ sia variabile in un intervallo aperto $N \subset \mathbb{R}$. Supponiamo che esista $\lambda_0 \in N$ per cui f_λ ammetta un punto fisso non iperbolico a $x_0 \in I$, cioè supponiamo che esistano $x_0 \in I$ e $\lambda_0 \in N$ tali che*

$$f_{\lambda_0}(x_0) = x_0; \quad f'_{\lambda_0}(x_0) \neq 1$$

Allora esistono intervalli aperti I_1 attorno a x_0 e N_1 attorno a λ_0 , e una funzione regolare $p(x) : N_1 \rightarrow I_1$ tali che $p(\lambda_0) = x_0$ e $f_\lambda(p(\lambda)) = p(\lambda)$. Inoltre $f_\lambda(x)$ non ha altri punti fissi in I_1 .

Osservazione 2.44

Il teorema precedente, così come quelli successivi, vale ovviamente per i punti periodici di periodo N semplicemente sostituendo f_λ con f_λ^N .

I due teoremi che seguono caratterizzano invece le biforcazioni nodo-sella e le biforcazioni che raddoppiano il periodo. Notiamo anzitutto che senza ridurre la generalità potremo sempre supporre nulla l'ascissa x_0 del punto di biforcazione.

Teorema 2.10 *Data la famiglia $f_\lambda(x)$ come nel teorema precedente, supponiamo quanto segue:*

1. $f_{\lambda_0}(0) = 0$;
2. $f'_{\lambda_0}(0) = 1$;
3. $f''_{\lambda_0}(0) \neq 0$;
4. $\left. \frac{\partial f}{\partial \lambda} \right|_{\lambda=\lambda_0}(0) \neq 0$.

Allora esiste un intervallo I_1 attorno a $x = 0$ e una funzione regolare $p : I_1 \rightarrow \mathbb{R}$ che soddisfa $p(0) = \lambda_0$ e tale che

$$f_{p(x)}(x) = x$$

Per di più $p'(0) = 0$ e $p''(0) \neq 0$.

Osservazione 2.45

I segni di $f''_{\lambda_0}(0)$ e $\left. \frac{\partial f}{\partial \lambda} \right|_{\lambda=\lambda_0}(0)$ determinano la "direzione" della biforcazione. Se ad esempio hanno segno opposto il diagramma di biforcazione appare come nella Fig.5.5.

Teorema 2.11 *Data la famiglia $f_\lambda(x)$ come nei teoremi precedenti, supponiamo ora quanto segue:*

1. $f_\lambda(0) = 0$ per tutti i λ in un intervallo attorno a $\lambda = \lambda_0$;
2. $f'_{\lambda_0}(0) = -1$;
3. $\left. \frac{\partial (f^2)'}{\partial \lambda} \right|_{\lambda=\lambda_0}(0) \neq 0$.

Allora esiste un intervallo I_1 attorno a P ed una funzione $p : I_1 \rightarrow \mathbb{R}$ tale che

$$f_{p(x)}(x) \neq x$$

ma

$$f_{p(x)}^2 = x$$

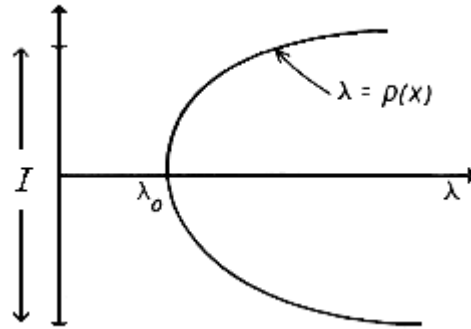


Figure 5.5: Il diagramma di biforcazione sia per la biforcazione nodo-sella che per quella a raddoppiamento di periodo. La curva $\lambda = p(x)$ dà i punti fissi per f_λ nel caos nodo-sella, e i punti di periodo 2 nel caso del raddoppiamento di periodo.

Esercizio 2.33

Identificare le biforcazioni e discutere le traiettorie di fase prima e dopo le biforcazioni per la mappa logistica $f := \mu x(1 - x)$ con $\mu = 3$.

Soluzione. Sappiamo già fin dall'Esempio 5.24 che $x = 0$ è un punto fisso repulsivo. Il secondo punto fisso è $x = p_\mu = \frac{\mu - 1}{\mu}$ (si veda l'Esempio 1.28) e quindi, per $\mu = 3$, $x = \frac{2}{3}$. Si ha poi $f'(\frac{2}{3}) = -1$, $f''(\frac{2}{3}) = -6$. Siamo quindi esattamente nelle condizioni del teorema precedente: il punto fisso $x = \frac{2}{3}$ dà ordine ad una biforcazione con raddoppiamento di periodo, cioè a due punti periodici di periodo 2, ambedue stabili.

3 Caos tramite infiniti raddoppiamenti di periodo

Ritorniamo alla domanda che ci siamo già posti all'inizio del Capitolo. la mappa quadratica $f_\mu(x) = \mu x(1 - x)$ genera una dinamica semplice per $\mu < 3$ e caotica per $\mu > 4$. Come fa a diventare caotica? Da dove nascono gli infiniti punti periodici che esistono per $\mu > 4$? Qui

cercheremo di darle una motivazione intuitiva, che sottoporremo poi alla verifica numerica tramite esperimenti al calcolatore. Tutto ciò che enunceremo può essere dimostrato rigorosamente, ma solo tramite conoscenze matematica più avanzate acquisibili solo in proseguendo il corso di studi.

L'idea che svilupperemo è la seguente: sappiamo che a $\mu = 3$ si incontra una prima biforcazione con raddoppiamento del periodo del punto fisso attrattivo a p_μ per $\mu > 3$. Le 2 orbite di periodo 2 sono punti fissi stabili $q_{1,2}$ di f_μ^2 . Al crescere di μ si forma una struttura autosimilare. Esiste un valore $\mu_1 > 3$ raggiunto il quale i due punti fissi di periodo 2 perdono la loro stabilità, perché $(f_{\mu_1}^2)'(q_{1,2}) = -1$. A $\mu = \mu_1$ ciascuno di questi due punti fissi genera una biforcazione con raddoppiamento di periodo. Dunque per $\mu > \mu_1$ ci saranno 4 orbite periodiche stabili di periodo 4 (4 punti fissi stabili per f_μ^4) finché non si incontra un valore μ_2 dove questi punti fissi perdono stabilità perché la derivata $(f_{\mu_2}^4)'$ tocca -1 . Il medesimo meccanismo di biforcazione genera $8 = 2^3$ orbite stabili di periodo 8, cioè 2^3 punti fissi stabili per $f_{\mu_3}^8$. Continuando con il medesimo meccanismo, ci aspettiamo di trovare una successione di valori $\mu_1, \mu_2, \dots, \mu_n, \dots$ tale che per $\mu > \mu_n$ ci siano 2^n orbite stabili di periodo 2^n , ciascuna di esse essendo un punto fisso stabile di $f_{\mu_n}^{2^n}$. Se quest successione ha un limite, denotato μ_∞ , ci aspettiamo che per $\mu > \mu_\infty$ la dinamica, dopo avere superato infinite biforcazioni che raddoppiano il periodo, diventi caotica.

Per $1 < \mu < 3$ f_μ ha un solo punto fisso attrattivo $p_\mu = (\mu - 1)\mu$, $0 < p_\mu < 1$. Si ha $f'(p_\mu) = 2 - \mu$. Definiamo *punto corrispondente* di p_μ rispetto alla mappa f_μ quel punto $\tilde{p}_\mu$ tale che $0 < \tilde{p}_\mu < p_\mu$ e $f_\mu(\tilde{p}_\mu) = p_\mu$.

Esercizio 3.34

Dimostrare che finché $f'(p_\mu) < 0$ il punto corrispondente $\tilde{p}_\mu$ esiste e vale $\frac{1}{\mu}$.

Soluzione Poniamo $\tilde{p}_\mu = x$. Basta allora risolvere l'equazione di secondo grado $\mu x(1 - x) = \frac{\mu - 1}{\mu}$. Le soluzioni sono $x_1 = \frac{\mu - 1}{\mu}$ e $x_2 = \frac{1}{\mu}$. x_1 ovviamente coincide col punto fisso p_μ ; se $\mu > 2$, cioè $f'(p_\mu) < 0$, $\frac{1}{\mu} < \frac{\mu - 1}{\mu} = p_\mu$, e quindi $\tilde{p}_\mu = \frac{1}{\mu}$.

Possiamo ora studiare la funzione f_μ^2 al variare del parametro μ e tracciarne il grafico. La cosa è ancora possibile perché si tratta di un polinomio di quarto grado di forma particolarmente

semplice:

$$f_\mu^2(x) = \mu(\mu x(1-x)[1-\mu x(1-x)]) = \mu^2 x(1-x)[1-\mu x(1-x)]$$

Esercizio 3.35

Studiare la funzione $g_\mu(x) := \mu_2 f_\mu^2(x)$ per $x \in [0, 1]$ e $2 < \mu < 4$.

Soluzione $g(x)_\mu$ è simmetrica attorno a $x = 1/2$. Poi: $g(p_\mu) = g(\tilde{p}_\mu) = \mu_2 p_\mu$. Per i valori fissati di x e μ si ha $g_\mu(x) = 0$ solo per $x = 0$ e $x = 1$, $\forall \mu$. Inoltre $g'_\mu(x) = (1-2x)[1-2\mu x(1-x)]$, $g''_\mu(x) = -2[1-2\mu x(1-x)] - 2\mu(1-2x)^2$. Dunque $g'_\mu(x) = 0$ per $x_0 = 1/2$ e $x_{1,2} = (1 \pm \sqrt{1-2/\mu})/2$. Si ha $g''_\mu(x_0) > 0$, $g''_\mu(x_{1,2}) < 0$. Quindi a $x = 1/2$ g_μ ha un minimo che vale $(4-\mu)/2^4$, mentre a $x = x_{1,2}$ ha due massimi. Nella Fig.5.5 si riporta il grafico di $g_\mu(x)$ per 6 valori di μ . Sebbene capovolto, esso assomiglia a quello di f_μ per differenti valori del parametro.

Infatti:

1. Dentro la scatola di lato $[\tilde{p}_\mu, p_\mu]$, cioè nella fossa, $f_\mu^2(x)$ ha un punto fisso al secondo estremo dell'intervallo $[\tilde{p}_\mu, p_\mu]$ perché $f_\mu^2(p_\mu) = p_\mu$ ed un unico punto critico entro l'intervallo¹.
2. Al crescere di μ la fossa si approfondisce finché non esce dalla scatola.

In altre parole, f_μ^2 si comporta sull'intervallo $[\tilde{p}_\mu, p_\mu]$ in modo molto simile a come la mappa originale f_μ si comporta sull'intervallo iniziale $[0, 1]$. In particolare, al crescere di μ ci aspettiamo come prima cosa di trovare un punto fisso di f_μ^2 (ciò un'orbita di periodo 2 di f_μ) in $[\tilde{p}_\mu, p_\mu]$. Successivamente anche questo punto fisso si boforcherà raddoppiando il periodo, esattamente come lo aveva fatto il punto fisso p_μ di f_μ . Si produrrà così un'orbita di periodo 4, o, equivalentemente, un punto periodico di periodo 4. Continuando il procedimento, troveremo scatole sempre più piccole dove i grafici di f_μ^4 , f_μ^8 saranno simili a quelli della funzione originale. Pertanto ci aspettiamo che f_μ subisca una successione di biforcazioni con raddoppiamento di periodo al crescere di μ .

¹Ricordiamo che la funzione f definita su un intervallo aperto $I \subset \mathbb{R}$ ed ivi derivabile ha un *punto critico* in $x_0 \in I$ se $f'(x_0) = 0$.

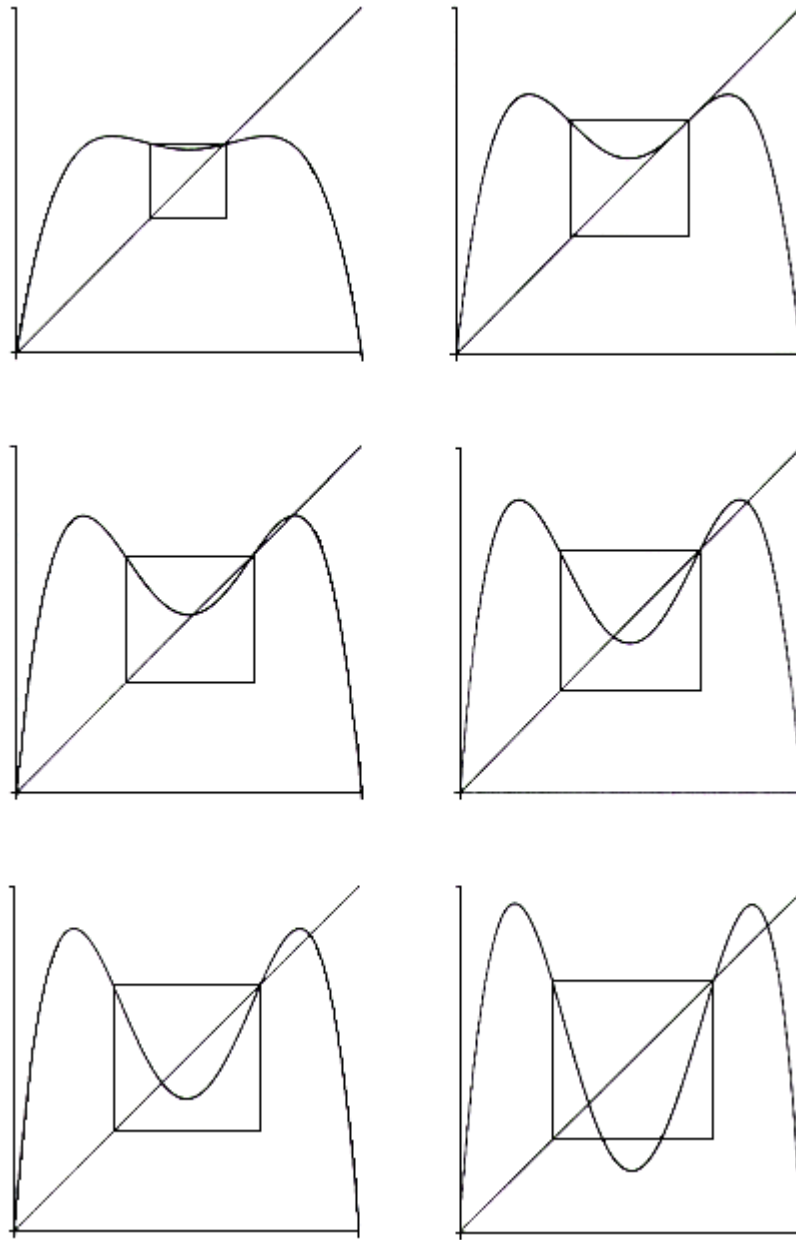


Figure 5.6: I grafici di $f_\mu^2(x)$ per $\mu = 2.5$, $\mu = 3.2$, $\mu = 3.4$, $\mu = 3.5$, $\mu = 3.8$.

Cerchiamo ora di rendere quantitative queste considerazioni. Poiché ci aspettiamo una struttura autosimilare o frattale, cioè una struttura che ha la medesima forma su qualsiasi scala la si voglia osservare, usiamo ancora lo strumento della "lente di ingrandimento". Per costruirne una conveniente, introduciamo il cosiddetto *operatore di rinormalizzazione* R . R è un operatore, nel senso che è una funzione che opera sulle funzioni: in altre parole il suo insieme di definizione è un insieme di funzioni. La sua azione è convertire certe funzioni definite su I in altre funzioni definite su I .

Sia ancora $\mu > 2$, e sia L_μ l'applicazione lineare che manda p_μ in 0 e $\tilde{p}_\mu$ in 1. Esplicitamente:

$$L_\mu(x) := \frac{x - p_\mu}{\tilde{p}_\mu - p_\mu}$$

L'inversa di $L_\mu(x)$ è evidentemente

$$L_\mu^{-1}(x) := (\tilde{p}_\mu - p_\mu)x + p_\mu$$

L'effetto di L_μ è quello di amplificare l'intervallo $[\tilde{p}_\mu, p_\mu]$ fino a $[0, 1]$ cambiandone l'orientamento.

Definiamo ora la rinormalizzazione di f_μ nel modo seguente:

$$(Rf_\mu)(x) = L_\mu \circ f_\mu^2 \circ L_\mu^{-1}(x) \quad (3.1)$$

La funzione rinormalizzata² Rf_μ è definita sull'intervallo $[0, 1]$ e condivide molte delle proprietà di f_μ stessa. Le enunciamo esplicitamente sotto forma di proposizione.

Proposizione 3.11 *L'effetto della trasformazione di rinormalizzazione è il seguente:*

$$(Rf_\mu)(x) = \mu^2 x(1-x)[1 - \mu x(1-x)] = f_\mu^2(x) \quad (3.2)$$

e pertanto valgono le seguenti proprietà

1. $(Rf_\mu)(0) = 0$ e $(Rf_\mu)(1) = 1$
2. $(Rf_\mu)'(1/2) = 0$ e $1/2$ è il solo punto critico di Rf_μ ;
3. $S(Rf_\mu) < 0$; qui S è la derivata di Schwarz.

²Rinormalizzare significa cambiare la normalizzazione, cioè la scala. Qui si rinormalizza in modo tale che l'intervallo di definizione della parte che interessa della mappa iterata sia sempre quello iniziale.

Osservazione 3.46

La derivata di Schwarz S di una funzione f derivabile almeno tre volte al punto x è definita così:

$$Sf(x) = \frac{f'''(x)}{f'(x)} - \frac{3}{2} \left(\frac{f''(x)}{f'(x)} \right)^2 \quad (3.3)$$

Questa nozione è utile per molti propositi. Qui enunciamo (dimostrazione in appendice) due delle sue proprietà più importanti:

Teorema 3.12

1. Sia P un polinomio, e siano le radici di $P'(x)$ reali e distinte. Allora $SP(x) > 0$;
2. Siano $Sf < 0$ e $Sg < 0$. Allora $S(f \circ g) < 0$.
3. Sia $-\infty \leq Sf < 0$. Supponiamo che f abbia n punti critici. Allora f ha al più $n + 2$ orbite periodiche attrattive.

Dimostrazione . Vedi Appendice.

Dimostrazione del Teorema 3.12. Si ha, inserendo le definizioni di L_μ e L_μ^{-1} ed eseguendo le composizioni:

$$(Rf_\mu)(x) = \mu^2 x(1-x)[1 - \mu x(1-x)] = f_\mu^2(x)$$

Questa da cui le affermazioni 1 e 2. Quanto alla 3, basta calcolare tre derivate. Ciò prova il Teorema.

Dunque la rinormalizzazione di f_μ trasforma punti periodici di periodo 2 di f_μ in punti fissi di Rf_μ . Inoltre, prima che μ arrivi a 4 il grafico di Rf_μ esce dal quadrato di lato unitario, cioè $[\tilde{p}_\mu, p_\mu]$.

Quindi, come abbiamo già notato, ci aspettiamo che Rf_μ subisca una biforcazione con raddoppiamento del periodo al crescere di μ . In effetti, finché Rf_μ continua ad ammettere un punto fisso con derivata negativa in un qualche punto $p_1(\mu)$, allora possiamo trovare il punto corrispondente $\tilde{p}_1(\mu)$ come in precedenza e definire una seconda rinormalizzazione. La mappa lineare stavolta porta $p_1(\mu)$ in 0 e $\tilde{p}_1(\mu)$ in 1 ed è pertanto una mappa differente, anche se definita in maniera completamente analoga a L . Pertanto vediamo che l'intero quadro si ripete e ritroviamo un'altra biforcazione che raddoppia il periodo, questa volta per

f_μ^2 . La continuazione di questo processo conduce a una successione di biforcazioni con raddoppiamento di periodo all'ulteriore scendere di μ . Pertanto ci aspettiamo che il diagramma di biforcazione di f_μ sia almeno così complicato come quello nella Fig.5.7.

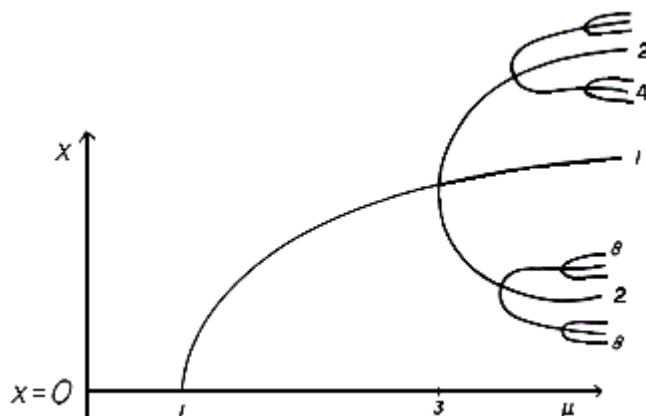


Figure 5.7: Diagramma di biforcazione di f_μ che mostra la successione dei raddoppiamenti di periodo. Gli interi rappresentano i periodi.

Il calcolatore ci permette di sottoporre queste considerazioni alla verifica sperimentale. Calcoliamo il *diagramma delle orbite* di f_μ . Si tratta di una figura che riporta l'andamento asintotico dell'orbita del punto $1/2$ per una varietà di valori di μ differenti fra 0 e 4. Per ogni μ , calcoliamo i primi 500 punti dell'iterazione di $\frac{1}{2}$ tramite f_μ , cioè i punti $f_\mu^n(1/2)$ per $n = 1, \dots, 500$, ma ne riportiamo solo gli ultimi 400 per eliminare gli effetti transienti. poi suddividiamo l'intervallo $0 \leq \mu \leq 4$ in 1500 parti uguali. La fig. 5.8 mostra il risultato dell'esperimento numerico: i valori di μ sono riportati in ascissa mentre in ordinata sono riportati i 400 punti dell'orbita corrispondente con punto iniziale $1/2$. Si noti come il diagramma di biforcazione sia immerso in quest'ultimo.

Osservazione 3.47

Se prendessimo un qualsiasi altro punto iniziale $0 < x_0 < 1$ otterremmo una figura di antura identica. Convien scegliere il punto $1/2$ perché si può dimostrare, a seguito del Teorema 3.12, che viene sempre attratto da un'orbita periodica attrattiva.

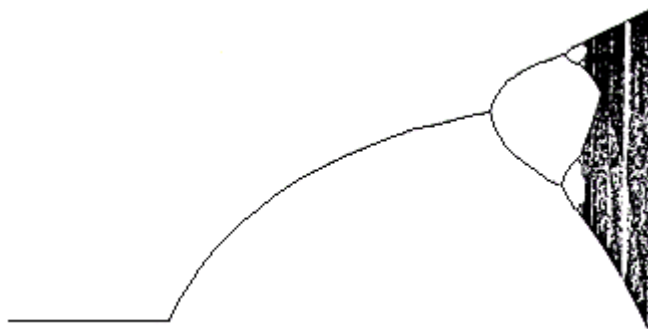


Figure 5.8: Il diagramma della orbite di f_μ per $0 \leq \mu \leq 4$. μ in ascissa.

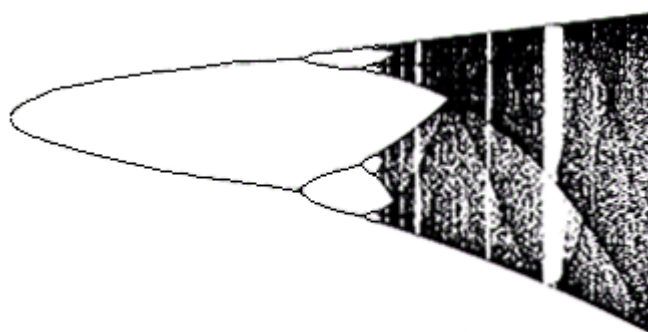


Figure 5.9: Il diagramma della orbite di f_μ per $3 \leq \mu \leq 4$

Vediamo chiaramente confermate dalla Fig.5.8 molte delle previsioni fin qui fatte. Per esempio per $0 \leq \mu \leq 1$ tutte le orbite convergono al punto fisso attrattivo $x = 0$. La convergenza è tanto più lenta quanto più $f'_\mu(0)$ si avvicina a 1, il che avviene per $\mu \rightarrow 1$. Questo rende conto del piccolo addensarsi di punti attorno a $(1, 9)$ nel diagramma. Per $1 \leq \mu < 3$ tutte le orbite sono attratte dal punto fisso $p_\mu = (\mu - 1)/\mu$, e lo si vede chiaramente nel diagramma delle orbite. Al crescere di μ si vede una successione di biforcazioni con raddoppiamento del periodo, il che conferma le nostre previsioni.

Si noti che per molti valori di μ oltre il regime di raddoppiamento del periodo sembra che l'orbita di $\frac{1}{2}$ riempi un intervallo. Naturalmente è molto difficile decidere se l'orbita è in realtà densa o se è attratta da un'orbita periodica di periodo molto elevato. Le precisione del calcolatore, limitata anche se molto elevata, non permette di distinguere fra questi due casi. Tuttavia, il diagramma delle orbite dà una certa prova sperimentale che molti dei valori di

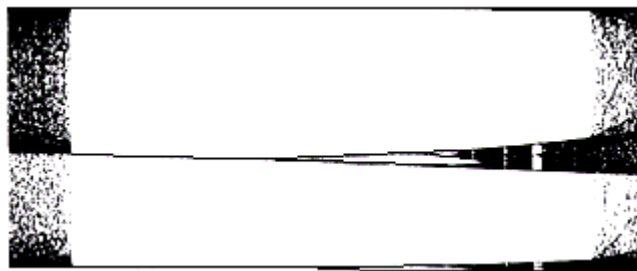


Figure 5.10: Il diagramma della orbite di f_μ per $3.825 \leq \mu \leq 3.86$

μ successivi a quelli in cui si ha raddoppiamento del periodo danno origine ad una dinamica caotica.

Questo aspetto lo si vede in modo più convincente nella Fig.5.9, dove abbiamo amplificato la porzione del diagramma delle orbite per $3 \leq \mu \leq 4$. Si noti che la successione delle biforcazioni con raddoppiamento del periodo è molto chiaramente visibile.

Il calcolatore ci permette di fare esperimenti sempre più interessanti, ingrandendo la finestra dove l'orbita di $\frac{1}{2}$ è attratta da un'orbita periodica di periodo n . Ingrandendo e ricalcolando questa porzione del diagramma delle orbite, si vede che si verifica sempre lo stesso fenomeno al crescere di μ : f_μ incontra un'altra successione di biforcazioni con raddoppiamento del periodo. Questo lo si vede bene nella Fig.5.10, dove si fa vedere la situazione del periodo 3 per $3.825 \leq \mu \leq 3.86$. Si noti che per molti valori di μ si vede solo un'orbita di periodo 3, che incontra una successione di di biforcazioni con raddoppiamento del periodo.

Capitolo 6

Dinamica discreta bidimensionale

1 Qualche richiamo di algebra lineare

Consideriamo il piano cartesiano $\mathbb{R}^2$, che possiamo sempre considerare come spazio vettoriale reale di dimensione 2 con la base canonica $\mathbf{e}_1 := (1, 0)$, $\mathbf{e}_2 := (0, 1)$. Su $\mathbb{R}^2$ si definisce anche il prodotto scalare euclideo $\langle \mathbf{x}, \mathbf{y} \rangle$ fra due vettori $\mathbf{x} = (x_1, x_2)$ e $\mathbf{y} = (y_1, y_2)$ tramite la formula $\langle \mathbf{x}, \mathbf{y} \rangle := x_1 y_1 + x_2 y_2$. Ricordiamo che un *endomorfismo*, o *trasformazione lineare* è un'applicazione $\mathbf{A} : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $\mathbf{x} \mapsto \mathbf{A}(\mathbf{x})$ che gode della proprietà di essere *lineare*. Un endomorfismo che sia anche invertibile si dice poi *automorfismo*. La proprietà di linearità significa che $\forall \mathbf{u}, \mathbf{v} \in \mathbb{R}^2$ e $\forall \alpha, \beta \in \mathbb{R}$ si ha

$$\mathbf{A}(\alpha \mathbf{u} + \beta \mathbf{v}) = \alpha \mathbf{A}(\mathbf{u}) + \beta \mathbf{A}(\mathbf{v}) \quad (1.1)$$

Essendo $\mathbf{A}(\mathbf{x})$ un vettore; esso avrà due componenti: $\mathbf{A}(\mathbf{x}) = (A_1, A_2) = A_1 \mathbf{e}_1 + A_2 \mathbf{e}_2$. Se $\mathbf{x} \in \mathbb{R}^2$, $\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2$, per la condizione di linearità (1.1) possiamo scrivere

$$\mathbf{A}(\mathbf{x}) = x_1 \mathbf{A}(\mathbf{e}_1) + x_2 \mathbf{A}(\mathbf{e}_2)$$

Ponendo: $\mathbf{A}(\mathbf{e}_1) = (a_{11}, a_{21})$, $\mathbf{A}(\mathbf{e}_2) = (a_{12}, a_{22})$ la formula precedente assume la forma

$$\mathbf{A}(\mathbf{x}) = x_1(a_{11}, a_{21}) + x_2(a_{12}, a_{22}) = (a_{11}x_1 + a_{12}x_2, a_{21}x_1 + a_{22}x_2) \quad (1.2)$$

Poiché $\mathbf{A}(\mathbf{x}) = (A_1, A_2)$ ne deduciamo

$$A_1 = a_{11}x_1 + a_{12}x_2 = \sum_{i=1}^2 a_{1i}x_i; \quad A_2 = a_{21}x_1 + a_{22}x_2 = \sum_{i=1}^2 a_{2i}x_i \quad (1.3)$$

Vale allora la seguente:

Proposizione 1.12 *Sia $(\mathbf{e}_1, \mathbf{e}_2)$ la base canonica in $\mathbb{R}^2$. Si consideri ogni vettore $x \in \mathbb{R}^2$ come un matrice colonna di ordine 2×1 i cui elementi sono le sue componenti cartesiane (x_1, x_2) , cioè le sue componenti rispetto alla base canonica: $\mathbf{x} := \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$. Allo stesso modo possiamo scrivere $\mathbf{A} := \begin{pmatrix} A_1 \\ A_2 \end{pmatrix}$. Allora $\mathbf{A}(\mathbf{x})$ agisce secondo la regola del prodotto di matrici riga per colonna nel modo seguente:*

$$\mathbf{A}(\mathbf{x}) = \begin{pmatrix} A_1 \\ A_2 \end{pmatrix} = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \quad (1.4)$$

dove

$$\mathcal{A} := \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$$

è la matrice 2×2 le cui colonne sono le componenti dei vettori $\mathbf{A}(\mathbf{e}_1)$ e $\mathbf{A}(\mathbf{e}_2)$ lungo la base canonica. In forma compatta la (1.4) si scrive

$$\mathbf{A}(\mathbf{x}) = \mathcal{A}\mathbf{x} \quad (1.5)$$

Dimostrazione. Si tratta di una verifica diretta.

Osservazione 1.48

1. La Proposizione precedente è un caso particolare di un teorema più generale. Esso afferma che, fissate due basi $E := (\mathbf{e}_1, \mathbf{e}_2)$ e $F := (\mathbf{f}_1, \mathbf{f}_2)$ in $\mathbb{R}^2$, $\forall x \in \mathbb{R}^2$ la matrice colonna 2×1 delle componenti di $\mathbf{A}(\mathbf{x})$ rispetto alla base F è data dal prodotto matriciale riga per colonna seguente:

$$\mathcal{A} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} := \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \quad (1.6)$$

dove: $\begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$ è la matrice delle componenti di $\mathbf{x}$ rispetto alla base E , cioè $\mathbf{x} = x_1\mathbf{e}_1 + x_2\mathbf{e}_2$, e la matrice $\mathcal{A}$ ha per colonne le componenti di $\mathbf{A}(\mathbf{e}_1)$ e $\mathbf{A}(\mathbf{e}_2)$ rispetto alla base F . Più precisamente, a_{ij} è la componente di $\mathbf{A}(\mathbf{e}_j)$ lungo $\mathbf{f}_i$, $i, j = 1, 2$.

Dunque fissate due basi qualsiasi l'endomorfismo A può essere rappresentato dalla matrice $\mathcal{A}$ tramite la (1.6). È immediato verificare che due endomorfismi diversi sono rappresentati da matrici diverse.

2. Nella Proposizione 1.12 entrambe le basi E e F coincidono con la base canonica e così si intenderà sempre nel seguito.

3. Viceversa, sia data una matrice 2×2 denotata $\mathcal{B}$:

$$\mathcal{B} = \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix}$$

Tramite moltiplicazione riga per colonna, come sopra, per un qualsiasi vettore $\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \in \mathbb{R}^2$ è facile da verificare che rimane definita una trasformazione lineare da $\mathbb{R}^2$ in sè:

$$B(\mathbf{x}) := \mathcal{B}\mathbf{x} = \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} b_{11}x_1 + b_{12}x_2 \\ b_{21}x_1 + b_{22}x_2 \end{pmatrix}$$

Abbiamo così dimostrato la seguente importante

Proposizione 1.13 *Ogni endomorfismo (o trasformazione) lineare da $\mathbb{R}^2$ in sè è rappresentato da una matrice 2×2 e viceversa, cioè:*

1. Se $\mathbf{x} \mapsto A(\mathbf{x})$ è un endomorfismo (trasformazione) lineare di $\mathbb{R}^2$, allora esso agisce così:

$$A(\mathbf{x}) = \mathcal{A}\mathbf{x} = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$$

dove gli elementi della matrice 2×2 $\mathcal{A}$ sono le componenti dei vettori $A(\mathbf{e}_1)$ e $A(\mathbf{e}_2)$ lungo i vettori della base canonica.:

2. Se

$$\mathcal{B} = \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix}$$

è una matrice 2×2 , allora essa definisce una trasformazione lineare $B(\mathbf{x})$ da $\mathbb{R}^2$ in sè la cui azione sui vettori $\mathbf{x}$ di componenti (x_1, x_2) sulla base canonica, ovvero $\mathbf{x} = x_1\mathbf{e}_1 + x_2\mathbf{e}_2$ cioè $\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$, è definita mediante la moltiplicazione riga per colonna:

$$B(\mathbf{x}) := \mathcal{B}\mathbf{x} = \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$$

Osservazione 1.49

1. Consideriamo la decomposizione canonica di $\mathbf{x} \in \mathbb{R}^2$, $\mathbf{x} = x_1\mathbf{e}_1 + x_2\mathbf{e}_2$. Poiché la base canonica è ortogonale, $\langle \mathbf{e}_i, \mathbf{e}_j \rangle = \delta_{ij}$, prendendo i prodotti scalari con $\mathbf{e}_1$ e $\mathbf{e}_2$ otteniamo subito: $x_1 = \langle \mathbf{x}, \mathbf{e}_1 \rangle$, $x_2 = \langle \mathbf{x}, \mathbf{e}_2 \rangle$. Considerando il caso particolare dei vettori

$\mathbf{A}(\mathbf{e}_1)$ e $\mathbf{A}(\mathbf{e}_2)$, le cui componenti lungo la base canonica sono (a_{11}, a_{21}) , (a_{12}, a_{22}) , rispettivamente, otteniamo la seguente formula per gli elementi della matrice $\mathcal{A}$:

$$\begin{aligned} a_{11} &= \langle \mathbf{e}_1, \mathbf{A}(\mathbf{e}_1) \rangle, & a_{12} &= \langle \mathbf{e}_1, \mathbf{A}(\mathbf{e}_2) \rangle \\ a_{21} &= \langle \mathbf{e}_2, \mathbf{A}(\mathbf{e}_1) \rangle, & a_{22} &= \langle \mathbf{e}_2, \mathbf{A}(\mathbf{e}_2) \rangle \end{aligned}$$

2. D'ora in poi identificheremo ogni trasformazione lineare A con la matrice 2×2 che la rappresenta sulla base canonica. Seguiremo poi il consueto abuso di notazione di indicare con A , e non più $\mathcal{A}$, la matrice stessa. Si noti che la trasformazione identità e la trasformazione nulla sono rappresentate dalle matrici I e O definite così:

$$I := C = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \text{ (matrice identica)} \quad O = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} \text{ (matrice nulla)}$$

3. La rappresentazione in termini di matrici quadrate delle trasformazioni lineari è la ragione per cui si definisce il prodotto delle matrici riga per colonna. Questo fatto viene messo ancor meglio in evidenza considerando la composizione di due trasformazioni lineari A e B . Vogliamo infatti fare vedere che se A e B sono due trasformazioni lineari rappresentate dalle omonime matrici 2×2 , allora $C = A \circ B$ è una trasformazione lineare rappresentata sulla base canonica dalla matrice $C = A \cdot B$, dove $A \cdot B$ denota il prodotto riga per colonna delle matrici A e B . Infatti:

$$\begin{aligned} B\mathbf{x} &= \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} b_{11}x_1 + b_{12}x_2 \\ b_{21}x_1 + b_{22}x_2 \end{pmatrix} \\ A \circ B\mathbf{x} &= A(B\mathbf{x}) = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \begin{pmatrix} b_{11}x_1 + b_{12}x_2 \\ b_{21}x_1 + b_{22}x_2 \end{pmatrix} = \\ &= \begin{pmatrix} a_{11}b_{11}x_1 + a_{12}b_{21}x_1 & a_{11}b_{12}x_2 + a_{12}b_{22}x_2 \\ a_{21}b_{11}x_1 + a_{22}b_{21}x_1 & a_{21}b_{12}x_2 + a_{22}b_{22}x_2 \end{pmatrix} = \\ &= \begin{pmatrix} a_{11}b_{11} + a_{12}b_{21} & a_{11}b_{12} + a_{12}b_{22} \\ a_{21}b_{11} + a_{22}b_{21} & a_{21}b_{12} + a_{22}b_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = C \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \end{aligned}$$

4. Abbiamo fin qui discusso le trasformazioni lineari e la loro rappresentazione in termini di matrici, per semplicità, nel caso bidimensionale del piano $\mathbb{R}^2$. È del tutto evidente che le considerazioni che precedono si applicano senza alcuna variazione al caso delle trasformazioni lineari in $\mathbb{R}^n$ che saranno rappresentate da matrici $n \times n$. Tuttavia nel seguito ci limiteremo a trattare le matrici 2×2 soprattutto per non appesantire inutilmente la notazione.

La rappresentazione in termini di matrici ci permette di caratterizzare immediatamente il nucleo $\mathcal{N}(A)$, e quindi le proprietà di invertibilità, di qualsiasi applicazione lineare A in $\mathbb{R}^2$. Ricordiamo infatti anzitutto che il *nucleo* di un'applicazione lineare A è la controimmagine del vettore nullo in $\mathbb{R}^2$ attraverso A , cioè:

$$\mathcal{N}(A) := \{\mathbf{x} \in \mathbb{R}^2 \mid A(\mathbf{x}) = 0\}$$

Vale allora il risultato ben noto dall'algebra lineare:

Proposizione 1.14 *Sia Ax una trasformazione lineare in $\mathbb{R}^2$ diversa dalla trasformazione identicamente nulla. Sia $\mathcal{A}$ la matrice che rappresenta A sulla base canonica. Allora vale la seguente alternativa:*

1. Il nucleo $\mathcal{N}(A) \neq \{0\}$ se e solo se $\det \mathcal{A} = 0$. In tal caso $\mathcal{N}(A)$ ha dimensione 1 ed è il sottospazio delle soluzioni del sistema lineare e omogeneo

$$\mathcal{A}\mathbf{x} = 0 \iff \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = 0 \iff \begin{cases} a_{11}x_1 + a_{12}x_2 = 0 \\ a_{21}x_1 + a_{22}x_2 = 0 \end{cases}$$

Questo sottospazio si ottiene fissando una qualunque soluzione non nulla del sistema precedente e prendendo lo spazio unidimensionale da essa generato.

2. A è iniettiva e suriettiva, cioè biunivoca su $\mathbb{R}^2$, se e solo se $\det \mathcal{A} \neq 0$. In tal caso la trasformazione inversa A^{-1} è rappresentata dalla matrice inversa:

$$\mathcal{A}^{-1} := \frac{1}{\det \mathcal{A}} \begin{pmatrix} a_{22} & -a_{21} \\ -a_{12} & a_{11} \end{pmatrix}$$

Ciò significa che $\forall \mathbf{u} \in \mathbb{R}^2$ l'equazione $A\mathbf{x} = \mathbf{u}$ ammette la soluzione unica $\mathbf{x} = A^{-1}\mathbf{u}$. In forma esplicita, questa affermazione significa che il sistema lineare non omogeneo nella due incognite x_1, x_2 :

$$A\mathbf{x} = \mathbf{u} \iff \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix} \iff \begin{cases} a_{11}x_1 + a_{12}x_2 = u_1 \\ a_{21}x_1 + a_{22}x_2 = u_2 \end{cases}$$

ammette la soluzione unica, data dalla formula di Cramer:

$$\begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \frac{1}{\det A} \begin{pmatrix} a_{22} & -a_{21} \\ -a_{12} & a_{11} \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \end{pmatrix} \iff \begin{cases} x_1 = \frac{a_{22}u_1 - a_{21}u_2}{\det A} \\ x_2 = \frac{-a_{12}u_1 + a_{11}u_2}{\det A} \end{cases}$$

Osservazione 1.50

1. Le operazioni di somma $C := A + B$ di due matrici 2×2 , e quella di moltiplicazione per un numero reale α , $D := \alpha A$, sono definite come sappiamo nel modo seguente:

$$C = \begin{pmatrix} c_{11} & c_{12} \\ c_{21} & c_{22} \end{pmatrix} = \begin{pmatrix} a_{11} + b_{11} & a_{12} + b_{12} \\ a_{21} + b_{21} & a_{22} + b_{22} \end{pmatrix}; \quad D = \begin{pmatrix} \alpha a_{11} & \alpha a_{12} \\ \alpha a_{21} & \alpha a_{22} \end{pmatrix}$$

Con queste operazioni le amtrici 2×2 formano uno spazio vettoriale.

Le matrici 2×2 formano poi come sappiamo un anello rispetto alle operazioni di somma elemento per elemento e moltiplicazione riga per colonna. Questo anello contiene dei divisori dello zero. In altre parole, la legge di annullamento del prodotto $AB = 0$ se e solo se $A = 0$ o $B = 0$ non vale. Ad esempio:

$$A := \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \Rightarrow A^2 = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$$

2. Le matrici invertibili sono tutte e sole quelle a determinante non nullo. Poiché $\det AB = \det A \cdot \det B$, anche AB e BA sono invertibili e quindi le matrici 2×2 invertibili formano gruppo (non abeliano) rispetto alla moltiplicazione riga per colonna, con elemento neutro I , la matrice identità. Questo gruppo è denotato $GL(2)$ (GL sta per generale lineare).
3. La proprietà $\det AB = \det A \cdot \det B$ (il determinante del prodotto vale il prodotto dei determinanti) può apparire misteriosa se dimostrata per via puramente algebrica. Il mistero scompare se si ricorre alla interpretazione geometrica del determinante: il determinante

$$d := \det \begin{pmatrix} x_1 & y_1 \\ x_2 & y_2 \end{pmatrix} = (x_1 y_2 - x_2 y_1)$$

è l'area *orientata* del parallelogramma costruito sui vettori di componenti (x_1, y_1) e (x_2, y_2) rispetto alla base canonica, cioè il parallelogramma che ha per vertici i 4 punti di coordinate $(0, 0)$, (x_1, y_1) , (x_2, y_2) e $(x_1 + y_1, x_2 + y_2)$. Si vede subito, infatti, che l'area di questo parallelogramma vale $x_1 y_2 - x_2 y_1$. Si conviene di prenderla orientata *positivamente* se tale numero è positivo, e *negativamente* se è negativo.

Esercizio 1.36

Dimostrare che l'area del parallelogramma così definito vale in effetti $x_1y_2 - x_2y_1$.

La proprietà $\det AB = \det A \cdot \det B$ esprime semplicemente il fatto che l'area del prodotto di due parallelogrammi indipendenti, cioè su piani differenti, è il prodotto delle aree.

2 Autovalori e autovettori. Iterazione di matrici

Data la matrice 2×2 A , poniamoci il problema di calcolarne le potenze successive A^n , $n = 2, 3, \dots$; $A^0 = I$, $A^1 = A$. Si vede subito che l'applicazione bovina della regola del prodotto riga per colonna dà origine a formule di grande complicazione perché alla potenza n si ottiene una matrice i cui elementi sono ciascuno la somma di 2^{n-1} termini.

Esercizio 2.37

Dimostrare quest'ultima affermazione.

Soluzione. Procediamo per induzione. L'affermazione è vera per $n = 1$. Assumiamola vera per $n - 1$:

$$A^{n-1} = \begin{pmatrix} a(n-1)_{11} & a(n-1)_{12} \\ a(n-1)_{21} & a(n-1)_{22} \end{pmatrix}$$

dove ognuno degli elementi $a(n-1)_{ij}$, $i, j = 1, 2$ è la somma di 2^{n-1} termini. Allora

$$\begin{aligned} A^n &= \begin{pmatrix} a(n-1)_{11} & a(n-1)_{12} \\ a(n-1)_{21} & a(n-1)_{22} \end{pmatrix} \begin{pmatrix} a_{11} & a_{12} \\ a_{12} & a_{22} \end{pmatrix} = \\ &= \begin{pmatrix} a(n-1)_{11}a_{11} + a(n-1)_{12}a_{21} & a(n-1)_{11}a_{12} + a(n-1)_{12}a_{22} \\ a(n-1)_{21}a_{11} + a(n-1)_{22}a_{21} & a(n-1)_{21}a_{12} + a(n-1)_{22}a_{22} \end{pmatrix} \end{aligned}$$

Quindi ogni elemento è esattamente la somma di $2^{n-1} \cdot 2 = 2^n$ termini.

Esistono tuttavia casi in cui il calcolo di A^n è immediato:

Esempio 2.38

1. A è diagonale: $A = \begin{pmatrix} \lambda & 0 \\ 0 & \mu \end{pmatrix}$, cioè tutti gli elementi fuori dalla diagonale principale sono nulli. In questo caso si calcola immediatamente:

$$A^n = \begin{pmatrix} \lambda^n & 0 \\ 0 & \mu^n \end{pmatrix}$$

2. A è triangolare, cioè tale che $a_{21} = 0$ (*triangolare superiore*) oppure $a_{12} = 0$ (*triangolare inferiore*). Consideriamo i casi particolari

$$A = \begin{pmatrix} 1 & \lambda \\ 0 & 1 \end{pmatrix}; \quad A = \begin{pmatrix} 1 & 0 \\ \lambda & 1 \end{pmatrix}$$

In tal caso si calcola

$$A^n = \begin{pmatrix} 1 & n\lambda \\ 0 & 1 \end{pmatrix} \text{ (superiore); } \quad A^n = \begin{pmatrix} 1 & 0 \\ n\lambda & 1 \end{pmatrix} \text{ (inferiore)} \quad (2.7)$$

3. A è una radice quadrata, cioè:

$$A = \begin{pmatrix} 0 & a \\ -a & 0 \end{pmatrix} \quad \text{oppure} \quad A = \begin{pmatrix} 0 & -a \\ a & 0 \end{pmatrix}$$

In tal caso si ha, per $n = 1, 2, \dots$

$$A^{2n} = (-1)^n a^{2n} I = a^{2n} \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}; \quad A^{2n+1} = (-1)^n a^{2n} A \quad (2.8)$$

4. A è della forma $A := \begin{pmatrix} \alpha & \beta \\ -\beta & \alpha \end{pmatrix}$ oppure $A := \begin{pmatrix} \alpha & \beta \\ \beta & \alpha \end{pmatrix}$ (si noti che l'esempio precedente è incluso in questo con $\alpha = 0$ e $\beta = a$). Allora vale la formula (nel secondo caso; formula analoga nel primo):

$$A^n = (\alpha^2 + \beta^2)^{n/2} \begin{pmatrix} \cos [n \arctg(\beta/\alpha)] & -\sin [n \arctg(\beta/\alpha)] \\ \sin [n \arctg(\beta/\alpha)] & \cos [n \arctg(\beta/\alpha)] \end{pmatrix} \quad (2.9)$$

I coefficienti di A^n possono poi essere riespressi come polinomi in α e β notando che

$$\cos \arctg(\beta/\alpha) = \frac{\alpha}{\sqrt{\alpha^2 + \beta^2}}, \quad \sin \arctg(\beta/\alpha) = \frac{\beta}{\sqrt{\alpha^2 + \beta^2}}$$

e facendo poi uso delle formule di moltiplicazione per i seni e i coseni.

Esercizio 2.38

Dimostrare le formule (2.7, 2.8, 2.9).

Soluzione. Per dimostrare la (2.7) procediamo per induzione. La formula è vera per $n = 1$. Ammettiamola vera per n , e dimostriamola per $n + 1$. Si ha, considerando il caso superiore:

$$A^{n+1} = A^n \cdot A = \begin{pmatrix} 1 & n\lambda \\ 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & \lambda \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & (n+1)\lambda \\ 0 & 1 \end{pmatrix}$$

il che prova l'asserzione per il caso triangolare superiore. Quello inferiore è completamente analogo.

Quanto alla (2.8), sia $J := \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$, oppure $J := \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$. Si trova subito che in entrambi i casi $J^2 = -I$ dove I , ricordiamo, è la matrice identità. Pertanto $J^4 = I$ e si vede quindi che le matrici J sono radici quarte dell'identità. Quindi $J^{2n} = (-1)^n I$, e $J^{2n+1} = J^{2n} J = (-1)^n I J = (-1)^n J$. Notando che possiamo scrivere $A = aJ$ otteniamo subito la (2.8). Passiamo ora alla (2.9). Moltiplicando e dividendo ogni elemento della matrice per $\sqrt{\alpha^2 + \beta^2}$ si trova

$$A = \sqrt{\alpha^2 + \beta^2} \begin{pmatrix} \frac{\alpha}{\sqrt{\alpha^2 + \beta^2}} & -\frac{\beta}{\sqrt{\alpha^2 + \beta^2}} \\ \frac{\beta}{\sqrt{\alpha^2 + \beta^2}} & \frac{\alpha}{\sqrt{\alpha^2 + \beta^2}} \end{pmatrix}$$

Consideriamo ora un triangolo rettangolo di cateti α e β . L'ipotenusa sarà $\sqrt{\alpha^2 + \beta^2}$, e per definizione di seno e coseno possiamo porre $\cos \theta = \frac{\alpha}{\sqrt{\alpha^2 + \beta^2}}$, $\sin \theta = \frac{\beta}{\sqrt{\alpha^2 + \beta^2}}$ da cui

$$A = \sqrt{\alpha^2 + \beta^2} \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} = \sqrt{\alpha^2 + \beta^2} R(\theta), \quad R(\theta) := \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$$

Ricordiamo ora che $R(\theta)$ è la matrice di rotazione degli assi coordinati di angolo θ rispetto ad un asse ortogonale al piano e passante per l'origine (ciò che rappresenta l'interpretazione geometrica dell'azione di A). Infatti, facciamola agire sui vettori della base canonica $\mathbf{e}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$, $\mathbf{e}_2 = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$. Si ha:

$$\begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \begin{pmatrix} \cos \theta \\ \sin \theta \end{pmatrix}, \quad \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \begin{pmatrix} -\sin \theta \\ \cos \theta \end{pmatrix}$$

Dunque il vettore $\mathbf{e}_1$ viene trasformato nel vettore di componenti $(\cos \theta, \sin \theta)$ e il vettore $\mathbf{e}_2$ nel vettore di componenti $(-\sin \theta, \cos \theta)$, il questa azione è per definizione una rotazione di angolo θ .

Proviamo ora per induzione il fatto molto intuitivo che facendo agire n volte $R(\theta)$, cioè facendo la potenza R^n , si ottiene una rotazione di angolo $n\theta$. Pertanto la potenza n -esima di R sarà la matrice di rotazione di angolo $n\theta$:

$$R(\theta)^n = \begin{pmatrix} \cos n\theta & -\sin n\theta \\ \sin n\theta & \cos n\theta \end{pmatrix} \quad (2.10)$$

Calcoliamo per cominciare $R(\theta)^2$. Si trova subito:

$$R(\theta)^2 = \begin{pmatrix} \cos^2 \theta - \sin^2 \theta & -\cos \theta \sin \theta - \sin \theta \cos \theta \\ 2 \sin \theta \cos \theta & -\sin^2 \theta + \cos^2 \theta \end{pmatrix} = \begin{pmatrix} \cos 2\theta & -\sin 2\theta \\ \sin 2\theta & \cos 2\theta \end{pmatrix}$$

Facciamo ora vedere, per completare l'induzione, che la (2.10) implica la validità della medesima formula con $n+1$ al posto di n . Si ha infatti:

$$\begin{aligned} R(\theta)^{n+1} &= \begin{pmatrix} \cos n\theta & -\sin n\theta \\ \sin n\theta & \cos n\theta \end{pmatrix} \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} = \\ &= \begin{pmatrix} \cos n\theta \cos \theta - \sin n\theta \sin \theta & -\cos n\theta \sin \theta - \sin n\theta \cos \theta \\ \sin n\theta \cos \theta + \cos n\theta \sin \theta & -\sin n\theta \cos \theta + \cos n\theta \sin \theta \end{pmatrix} \\ &= \begin{pmatrix} \cos (n+1)\theta & -\sin (n+1)\theta \\ \sin (n+1)\theta & \cos (n+1)\theta \end{pmatrix} \end{aligned}$$

Quindi

$$A^n = (\alpha^2 + \beta^2)^{n/2} R(\theta)^n = (\alpha^2 + \beta^2)^{n/2} \begin{pmatrix} \cos [n \arctg(\beta/\alpha)] & -\sin [n \arctg(\beta/\alpha)] \\ \sin [n \arctg(\beta/\alpha)] & \cos [n \arctg(\beta/\alpha)] \end{pmatrix}$$

perché $\theta = \arctg(\beta/\alpha)$.

Gli esempi precedenti sono solo apparentemente particolari. Vedremo infatti che, a patto di scegliere la base giusta, e di ammettere che λ e μ possano assumere anche valori complessi, *tutte* le matrici 2×2 assumono la forma diagonale o triangolare. In altre parole, la applicazione lineare è rappresentata da una matrice non diagonale o non triangolare perché se ne cerca la rappresentazione sotto forma di matrice nella base canonica; se si trova invece la base giusta, essa assume la forma voluta. Per trovare la base giusta, dobbiamo introdurre alcune nozioni importanti e generali.

Definizione 2.33 *Sia A una matrice 2×2 . Allora:*

1. *Il polinomio di secondo grado in λ $p(\lambda)$ definito come*

$$p(\lambda) := \det(A - \lambda I) = \det \begin{pmatrix} a_{11} - \lambda & a_{12} \\ a_{21} & a_{22} - \lambda \end{pmatrix} \quad (2.11)$$

si dice polinomio caratteristico della matrice A . Esplicitamente:

$$p(\lambda) = \lambda^2 - (a_{11} + a_{22})\lambda + (a_{11}a_{22} - a_{12}a_{21}) \quad (2.12)$$

2. Ogni radice (reale o complessa) del polinomio caratteristico, cioè ogni soluzione dell'equazione caratteristica della matrice A :

$$p(\lambda) = 0 \quad (2.13)$$

si dice *autovalore* di A , di molteplicità pari alla sua molteplicità come radice dell'equazione caratteristica.

3. Sia λ un autovalore di A . Ogni vettore $\mathbf{v} \neq 0$ tale che

$$A\mathbf{v} = \lambda\mathbf{v} \quad (2.14)$$

si dice *autovettore* di A corrispondente all'autovalore λ .

Osservazione 2.51

1. Si noti che su un autovettore $\mathbf{v}$ la trasformazione lineare A agisce semplicemente come moltiplicazione per il numero λ (l'autovalore): $A\mathbf{v} = \lambda\mathbf{v}$. Quindi $A^n\mathbf{v} = \lambda^n\mathbf{v} \forall n$.
2. Noto l'autovalore λ , gli autovettori corrispondenti sono i vettori $\mathbf{v}$ tali che $(A - \lambda I)\mathbf{v} = 0$, cioè gli elementi del nucleo $\mathcal{N}(A - \lambda I)$ della trasformazione lineare $A - \lambda I$. Essi si trovano risolvendo l'equazione (2.14), cioè $(A - \lambda I)\mathbf{v} = 0$. Scritta esplicitamente per componenti, dato che $\mathbf{v} = (v_1, v_2)$ essa equivale al sistema lineare omogeneo

$$\begin{cases} (a_{11} - \lambda)v_1 + a_{12}v_2 = 0 \\ a_{21}v_1 + (a_{22} - \lambda)v_2 = 0 \end{cases} \quad (2.15)$$

Questo sistema ha soluzioni diverse da quella banale perché il determinante dei coefficienti è per ipotesi nullo dato che λ è autovalore. Naturalmente due qualunque di questi autovettori sono linearmente dipendenti perché uno è multiplo dell'altro. Più precisamente, notiamo che l'insieme degli autovettori corrispondente ad un medesimo autovalore λ , è uno spazio vettoriale, di dimensione non superiore a 2 perché è un sottospazio di $\mathbb{R}^2$, detto *autospatio* di λ . Infatti se $\mathbf{u}_1$ e $\mathbf{u}_2$ sono autovettori di A corrispondenti all'autovalore λ , $A\mathbf{u}_1 = \lambda\mathbf{u}_1$, $A\mathbf{u}_2 = \lambda\mathbf{u}_2$, si ha per la linearità di A :

$$A(\alpha\mathbf{u}_1 + \beta\mathbf{u}_2) = \alpha A\mathbf{u}_1 + \beta A\mathbf{u}_2 = \alpha\lambda\mathbf{u}_1 + \beta\lambda\mathbf{u}_2 = \lambda(\alpha\mathbf{u}_1 + \beta\mathbf{u}_2) \forall (\alpha, \beta) \in \mathbb{R}$$

e pertanto anche $\lambda(\alpha\mathbf{u}_1 + \beta\mathbf{u}_2)$ è autovettore corrispondente all'autovalore λ .

L'esempio più semplice è la matrice identità $A = I$. Essa ammette il solo autovalore 1 , e l'autospazio è tutto $\mathbb{R}^2$.

3. Si noti che ogni autovettore è definito a meno di una costante moltiplicativa (come lo sono, del resto, tutte le soluzioni non banali di un sistema lineare omogeneo): se $\mathbf{v}$ è autovettore di A , $A\mathbf{v} = \lambda\mathbf{v}$, allora anche $\mathbf{w} := \alpha\mathbf{v}$ lo è $\forall \alpha \neq 0$. Naturalmente tutti questi vettori sono linearmente dipendenti fra loro.
4. Poiché la somma delle radici del polinomio caratteristico vale il coefficiente del termine di primo grado, e il loro prodotto quello di grado 0, per la (2.12) si ha $\lambda_1 + \lambda_2 = a_{11} + a_{22}$, $\lambda_1 \cdot \lambda_2 = a_{11} \cdot a_{22} - a_{12} \cdot a_{21} = \det(A)$. Se definiamo *traccia* della matrice A , denotata $\text{Tr}A$, la somma dei suoi elementi diagonali, possiamo concludere che $\lambda_1 + \lambda_2 = \text{Tr}A$. Pertanto l'equazione caratteristica $p(\lambda) = 0$ può essere riscritta così:

$$\lambda^2 - \text{Tr}A\lambda + \det(A) = 0 \quad (2.16)$$

5. Dato che $\det(A) = \lambda_1 \cdot \lambda_2$, la condizione necessaria e sufficiente "algebrica" affinché A sia invertibile è che i suoi autovalori siano tutti diversi da 0. La stessa cosa si può vedere per via "geometrica" osservando che $\lambda = 0$ è autovalore se e solo se $\exists \mathbf{u} \neq 0$ tale che $A\mathbf{u} = 0$ il che equivale a dire che A non è invertibile.

Esercizio 2.39

Trovare autovalori e autovettori delle seguenti matrici 2×2 :

$$D := \begin{pmatrix} \mu & 0 \\ 0 & \nu \end{pmatrix}; \quad A := \begin{pmatrix} \alpha & \beta \\ -\beta & \alpha \end{pmatrix}; \quad \Lambda = \begin{pmatrix} 1 & a \\ 0 & 1 \end{pmatrix} \quad (2.17)$$

Soluzione. Il caso della matrice diagonale D è banale: si ha subito $\lambda_1 = \mu$, $\mathbf{v}_1 = \mathbf{e}_1$, $\lambda_2 = \nu$, $\mathbf{v}_2 = \mathbf{e}_2$. Trattiamo ora A . L'equazione caratteristica (2.13) qui diventa, tenendo conto della definizione (2.11):

$$\lambda^2 - 2\alpha\lambda + \alpha^2 + \beta^2 = 0$$

che ammette le due soluzioni complesse coniugate $\lambda_{1,2} = \alpha \pm i\beta$. A ammette quindi i due autovalori $\lambda_1 = \alpha + i\beta$, $\lambda_2 = \alpha - i\beta$, distinti e quindi di molteplicità 1 (equivalentemente, *semplici* (un autovalore di molteplicità 1 si dice semplice) se $\beta \neq 0$). Calcoliamo gli autovettori risolvendo il sistema lineare omogeneo (2.15) cominciando da $\lambda = \lambda_1 = \alpha + i\beta$. Il sistema è:

$$\begin{cases} (\alpha - \lambda_1)v_1 + \beta v_2 = 0 \\ -\beta v_1 + (\alpha - \lambda_1)v_2 = 0 \end{cases} \quad \text{ovvero} \quad \begin{cases} -i\beta v_1 + \beta v_2 = 0 \\ -\beta v_1 - i\beta v_2 = 0 \end{cases}$$

Si noti che le due equazioni dell'ultimo sistema sono in realtà una sola, così come deve essere dato che il determinante del sistema è nullo. Per risolvere il sistema omogeneo, fissando ad esempio $v_1 = 1$, otteniamo $v_2 = i$. Quindi un autovettore $\mathbf{v}_1$ corrispondente a λ_1 è $\mathbf{v}_1 = \begin{pmatrix} 1 \\ i \end{pmatrix} = \mathbf{e}_1 + i\mathbf{e}_2$, e ogni altro si ottiene da questo tramite moltiplicazione per una costante arbitraria. Analogamente si ottiene $\mathbf{v}_2 = \begin{pmatrix} 1 \\ -i \end{pmatrix} = \mathbf{e}_1 - i\mathbf{e}_2$ per l'autovettore corrispondente a λ_2 . Si noti che questi autovettori sono complessi, e quindi non appartengono a $\mathbb{R}^2$. Il modo più semplice e naturale di risolvere questa contraddizione apparente è quello di fare agire A su $\mathbb{C}^2$ invece che su $\mathbb{R}^2$. $\mathbb{C}^2$ con le operazioni di somma e moltiplicazione per uno scalare (complesso) è infatti uno spazio vettoriale al pari di $\mathbb{R}^2$, e la matrice A , che possiamo assumere ad elementi complessi, agisce sempre tramite moltiplicazione riga per colonna, stavolta fra numeri complessi.

Passiamo ora a A . L'equazione caratteristica sarà

$$(1 - \lambda)^2 = 0$$

che ammette la radice doppia $\lambda = 1$. Quindi c'è un solo autovalore di molteplicità 2. Per calcolare gli autovettori, scriviamo anche qui il sistema omogeneo $(A - \lambda I)\mathbf{v} = 0$, che ora assume evidentemente la forma

$$\begin{cases} av_2 = 0 \\ 0 = 0 \end{cases}$$

le cui soluzioni sono $v_2 = 0$ e $a \neq 0$ e $v_1 = v$ con $v \in \mathbb{R}$ arbitrario. Poiché $(v, 0) = v(1, 0)$ possiamo prendere $\mathbf{e}_1 = (1, 0)$ come autovettore e tutti gli altri si ottengono da questo tramite moltiplicazione per lo scalare v . In questo caso quindi nonostante ci siano due autovalori (ricordiamo che ogni autovalore va contato un numero di volte pari alla sua molteplicità) l'autospazio generato dagli autovettori ha dimensione 1.

I casi delle matrici D , A e Λ sono solo apparentemente particolari. In realtà essi essenzialmente esauriscono le matrici 2×2 come ora vedremo.

Definizione 2.34 Due trasformazioni lineari A e B in $\mathbb{R}^2$ (identificate al solito con le matrici A e B) si dicono simili o linearmente coniugate se esiste una trasformazione lineare invertibile Q (detta similitudine) tale che

$$A = QBQ^{-1} \tag{2.18}$$

Osservazione 2.52

Se due trasformazioni A e B sono coniugate dalla similitudine Q , la stessa proprietà vale per ogni loro potenza A^n e B^n : $A^n = QB^nQ^{-1}$. Infatti, applicando n volte la (2.18) si ottiene:

$$A^n = [QBQ^{-1}]^n = [QBQ^{-1}QBQ^{-1} \cdots QBQ^{-1}] = QB^nQ^{-1}$$

Teorema 2.13 *Siano A e B due trasformazioni lineari simili. Allora:*

1. *A e B hanno gli stessi autovalori;*

2. *Sia $L := \begin{pmatrix} l_{11} & l_{12} \\ l_{21} & l_{22} \end{pmatrix}$ una matrice 2×2 arbitraria. Allora esiste una trasformazione di similitudine Q tale che QLQ^{-1} assume una delle tre forme D , A o Λ della (2.17).*

Dimostrazione. L'equazione caratteristica per A è $\det(A - \lambda I) = 0$, e quella di B $\det(B - \lambda I) = 0$. Ora si ha

$$\det(B - \lambda I) = \det(QAQ^{-1} - \lambda QQ^{-1}) = \det[Q(A - \lambda I)Q^{-1}] = \det Q \det(A - \lambda I) \det Q^{-1}$$

dato che $QQ^{-1} = I$ e che il determinante del prodotto vale il prodotto dei determinanti. Ora:

$$1 = \det I = \det(QQ^{-1}) = \det Q \det Q^{-1} \implies \det Q = 1/\det Q^{-1} \neq 0.$$

Quindi le due equazioni $\det(A - \lambda I) = 0$ e $\det(B - \lambda I) = 0$ hanno le stesse radici. Questo prova l'affermazione 1.

Per provare l'affermazione 2 costruiamo esplicitamente la matrice di coniugazione Q . Cominciamo con l'assumere che L abbia autovalori complessi $\alpha \pm i\beta$. Ricordiamo l'osservazione del Capitolo 1 che gli autovalori di L , se complessi, devono essere l'uno il coniugato dell'altro perché l'equazione caratteristica ha coefficienti reali. Procedendo esattamente come nella soluzione dell'esercizio precedente si trova subito che in tal caso esistono due autovettori linearmente indipendenti $\mathbf{w}_1$ e $\mathbf{w}_2$ le cui componenti sono a priori complesse. Separando parti reali e immaginarie scriviamo allora $\mathbf{w}_1 = \mathbf{v}_1 + i\mathbf{v}_2$ dove ora $\mathbf{v}_1$ e $\mathbf{v}_2$ sono vettori reali. Poiché $L\mathbf{w} = (\alpha + i\beta)\mathbf{w}$ ne segue, separando ancora parte reale e parte immaginaria, cioè scrivendo $L\mathbf{w} = L\mathbf{v}_1 + iL\mathbf{v}_2$, $(\alpha + i\beta)\mathbf{w} = (\alpha + i\beta)(\mathbf{v}_1 + i\mathbf{v}_2) = (\alpha\mathbf{v}_1 - \beta\mathbf{v}_2) + i(\beta\mathbf{v}_1 + \alpha\mathbf{v}_2)$:

$$\begin{aligned} L\mathbf{v}_1 &= \alpha\mathbf{v}_1 - \beta\mathbf{v}_2; \\ L\mathbf{v}_2 &= \beta\mathbf{v}_1 + \alpha\mathbf{v}_2 \end{aligned}$$

ovvero, scrivendo esplicitamente le componenti:

$$\begin{pmatrix} l_{11}v_{11} + l_{12}v_{12} \\ l_{21}v_{11} + l_{22}v_{12} \end{pmatrix} = \begin{pmatrix} \alpha v_{11} - \beta v_{21} \\ \alpha v_{12} - \beta v_{22} \end{pmatrix}; \quad \begin{pmatrix} l_{11}v_{21} + l_{12}v_{22} \\ l_{21}v_{21} + l_{22}v_{22} \end{pmatrix} = \begin{pmatrix} \beta v_{11} + \alpha v_{21} \\ \beta v_{12} + \alpha v_{22} \end{pmatrix} \quad (2.19)$$

Sia ora Q la matrice 2×2 le cui colonne sono i vettori $\mathbf{v}_1$ e $\mathbf{v}_2$, o, equivalentemente, la matrice Q tale che $Q\mathbf{e}_j = \mathbf{v}_j$, cioè $\mathbf{e}_j = Q^{-1}\mathbf{v}_j$:

$$Q := \begin{pmatrix} v_{11} & v_{21} \\ v_{12} & v_{22} \end{pmatrix} \quad (2.20)$$

Q è invertibile perché i vettori $\mathbf{v}_1$ e $\mathbf{v}_2$ sono linearmente indipendenti. Infatti:

$$\det Q = 0 \iff v_{11}v_{22} - v_{21}v_{12} = 0 \iff \frac{v_{11}}{v_{12}} = \frac{v_{21}}{v_{22}}$$

il che equivale alla lineare dipendenza di $\mathbf{v}_1$ e $\mathbf{v}_2$.

Allora poniamo $Q^{-1} := \begin{pmatrix} m_{11} & m_{12} \\ m_{21} & m_{22} \end{pmatrix}$ e calcoliamo direttamente:

$$\begin{aligned} Q^{-1}LQ &= \begin{pmatrix} m_{11} & m_{12} \\ m_{21} & m_{22} \end{pmatrix} \begin{pmatrix} l_{11} & l_{12} \\ l_{21} & l_{22} \end{pmatrix} \begin{pmatrix} v_{11} & v_{21} \\ v_{12} & v_{22} \end{pmatrix} \\ &= \begin{pmatrix} m_{11} & m_{12} \\ m_{21} & m_{22} \end{pmatrix} \begin{pmatrix} l_{11}v_{11} + l_{12}v_{12} & l_{11}v_{21} + l_{12}v_{22} \\ l_{21}v_{11} + l_{22}v_{12} & l_{21}v_{21} + l_{22}v_{22} \end{pmatrix} \\ &= \begin{pmatrix} m_{11} & m_{12} \\ m_{21} & m_{22} \end{pmatrix} \begin{pmatrix} \alpha v_{11} - \beta v_{21} & \beta v_{11} + \alpha v_{21} \\ \alpha v_{12} - \beta v_{22} & \beta v_{12} + \alpha v_{22} \end{pmatrix} = \begin{pmatrix} \alpha & \beta \\ -\beta & \alpha \end{pmatrix} \end{aligned}$$

dove la prima uguaglianza è conseguenza del calcolo del prodotto riga per colonna LQ , la seconda dell'applicazione della (2.19), e la terza del fatto che, essendo $L^{-1}L = I$, si ha, applicando ancora la regola del prodotto riga per colonna,

$$\begin{aligned} m_{11}v_{11} + m_{12}v_{21} &= 1 & m_{11}v_{21} + m_{12}v_{22} &= 0 \\ m_{21}v_{11} + m_{22}v_{21} &= 0 & m_{21}v_{21} + m_{22}v_{22} &= 1 \end{aligned} \quad (2.21)$$

Passiamo ora al caso in cui L ha autovalori reali, che possono essere coincidenti o distinti. Supponiamo anzitutto che siano distinti $\lambda_1 \neq \lambda_2$, e quindi semplici. Anche qui, risolvendo i due sistemi omogenei $(L - \lambda_1)\mathbf{v} = 0$ e $(L - \lambda_2)\mathbf{v} = 0$, troviamo che esistono due autovettori linearmente indipendenti $\mathbf{v}_1$ in corrispondenza a λ_1 e $\mathbf{v}_2$ in corrispondenza a λ_2 . Definiamo la matrice Q esattamente come sopra, formula (2.20), e notiamo che poiché stavolta $L\mathbf{v}_1 = \lambda_1\mathbf{v}_1, L\mathbf{v}_2 = \lambda_2\mathbf{v}_2$ le (2.19) diventano

$$\begin{pmatrix} l_{11}v_{11} + l_{12}v_{12} \\ l_{21}v_{11} + l_{22}v_{12} \end{pmatrix} = \lambda_1 \begin{pmatrix} v_{11} \\ v_{12} \end{pmatrix}; \quad \begin{pmatrix} l_{11}v_{21} + l_{12}v_{22} \\ l_{21}v_{21} + l_{22}v_{22} \end{pmatrix} = \lambda_2 \begin{pmatrix} v_{21} \\ v_{22} \end{pmatrix} \quad (2.22)$$

Pertanto, procedendo come sopra

$$\begin{aligned}
 Q^{-1}LQ &= \begin{pmatrix} m_{11} & m_{12} \\ m_{21} & m_{22} \end{pmatrix} \begin{pmatrix} l_{11} & l_{12} \\ l_{21} & l_{22} \end{pmatrix} \begin{pmatrix} v_{11} & v_{21} \\ v_{12} & v_{22} \end{pmatrix} \\
 &= \begin{pmatrix} m_{11} & m_{12} \\ m_{21} & m_{22} \end{pmatrix} \begin{pmatrix} l_{11}v_{11} + l_{12}v_{12} & l_{11}v_{21} + l_{12}v_{22} \\ l_{21}v_{11} + l_{22}v_{12} & l_{21}v_{21} + l_{22}v_{22} \end{pmatrix} \\
 &= \begin{pmatrix} m_{11} & m_{12} \\ m_{21} & m_{22} \end{pmatrix} \begin{pmatrix} \lambda_1 v_{11} & \lambda_2 v_{21} \\ \lambda_1 v_{12} & \lambda_2 v_{22} \end{pmatrix} = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix}
 \end{aligned}$$

ancora in virtù delle (2.21). La conclusione

eh che $QLQ^{-1} = D$ dove D è una matrice diagonale della forma voluta.

Rimane da trattare il caso dei due autovalori coincidenti: $\lambda_1 = \lambda_2 = \lambda$. Distinguiamo due sottocasi. Il primo è quello in cui a λ corrispondono due autovettori linearmente indipendenti $\mathbf{v}_1$ e $\mathbf{v}_2$, o, equivalentemente, il sistema lineare e omogeneo $(L - \lambda I)\mathbf{v} = 0$ ammette due soluzioni linearmente indipendenti $\mathbf{v}_1$ e $\mathbf{v}_2$. In tal caso possiamo applicare senza alcuna variazione il ragionamento del caso precedente e concludiamo che

$$QLQ^{-1} = D, \quad D := \begin{pmatrix} \lambda & 0 \\ 0 & \lambda \end{pmatrix}$$

Rimane da esaminare il caso in cui il sistema lineare e omogeneo $(L - \lambda I)\mathbf{v} = 0$ ammette una sola soluzione linearmente indipendente $\mathbf{v}_1$. In tal caso esiste $\mathbf{v}_2 \neq 0$ tale che $(L - \lambda I)\mathbf{v}_2 \neq 0$. Tuttavia dovrà risultare $(L - \lambda I)^2\mathbf{v}_2 = 0$. Infatti se $\mathbf{v}_1$ e $\mathbf{v}_2$ sono linearmente indipendenti, essi formano una base in $\mathbb{R}^2$ e si può scrivere $(L - \lambda I)\mathbf{v}_2 = a\mathbf{v}_1 + b\mathbf{v}_2$. Pertanto, dato che $(L - \lambda I)\mathbf{v}_1 = 0$:

$$(L - \lambda I)(L - \lambda I)\mathbf{v}_2 = \lambda(L - \lambda I)\mathbf{v}_2 = b(L - \lambda I)\mathbf{v}_2 \implies (L - \lambda I - b)(L - \lambda I)\mathbf{v}_2 = 0.$$

D'altra parte λ è il solo autovalore di L ; quindi $b = 0$ e $(L - \lambda I)^2\mathbf{v}_2 = 0$.

Definiamo ora, in analogia ai casi precedenti, la matrice Q come la matrice formata mettendo in colonna i vettori $\mathbf{w}_1$ e $\mathbf{w}_2$, che sono linearmente indipendenti:

$$Q = \begin{pmatrix} w_{11} & w_{21} \\ w_{12} & w_{22} \end{pmatrix}$$

Le (2.19) qui diventano

$$\begin{pmatrix} l_{11}w_{11} + l_{12}w_{12} \\ l_{21}w_{11} + l_{22}w_{12} \end{pmatrix} = \begin{pmatrix} \lambda w_{11} + w_{21} \\ \lambda w_{12} + w_{22} \end{pmatrix}; \quad \begin{pmatrix} l_{11}w_{21} + l_{12}w_{22} \\ l_{21}w_{21} + l_{22}w_{22} \end{pmatrix} = \lambda \begin{pmatrix} w_{21} \\ w_{22} \end{pmatrix} \quad (2.23)$$

Ora ricordiamo che per definizione $Q\mathbf{e}_j = \mathbf{w}_j$ ovvero $\mathbf{e}_j = Q^{-1}\mathbf{w}_j$. Pertanto:

$$Q^{-1}LQ = Q^{-1}(\lambda\mathbf{w}_1 + \mathbf{w}_2, \lambda\mathbf{w}_2) = (\lambda\mathbf{e}_1 + \mathbf{e}_2, \lambda\mathbf{e}_2) = \begin{pmatrix} \lambda & 0 \\ 1 & \lambda \end{pmatrix}$$

che è la forma triangolare voluta. Questo conclude la dimostrazione del teorema.

Osservazione 2.53

1. Dunque una trasformazione di similitudine lascia invariati gli autovalori. Gli autovettori vengono invece trasformati secondo la matrice di coniugazione Q . Supponiamo infatti che B abbia un autovettore $\mathbf{u}$: $B\mathbf{u} = \lambda\mathbf{u}$. Allora la (2.18) può essere riscritta, dopo avere moltiplicato a destra ambo i membri per Q , come $AQ = QB$. Otteniamo subito allora

$$AQ\mathbf{u} = QB\mathbf{u} = Q\lambda\mathbf{u} = \lambda Q\mathbf{u}$$

e quindi $Q\mathbf{u}$ è autovettore di A corrispondente all'autovalore λ . In modo analogo se, inversamente, $\mathbf{v}$ è autovettore di A corrispondente all'autovalore μ allora $Q^{-1}\mathbf{v}$ è autovettore di B corrispondente al medesimo autovalore. In particolare, se A ammette due autovettori linearmente indipendenti $\mathbf{u}_1$ e $\mathbf{u}_2$ lo stesso vale per B che, poiché Q è invertibile, ammette i due autovettori linearmente indipendenti $Q^{-1}\mathbf{u}_1$ e $Q^{-1}\mathbf{u}_2$. In altre parole, i vettori $Q^{-1}\mathbf{u}_1$ e $Q^{-1}\mathbf{u}_2$ formano una base in $\mathbb{C}^2$ al pari di $\mathbf{u}_1$ e $\mathbf{u}_2$. Infatti se Q è biunivoca Q e Q^{-1} trasformano basi in basi. Sappiamo infatti che condizione necessaria e sufficiente affinché due vettori $\mathbf{u}_1$ e $\mathbf{u}_2$ formino una base (in $\mathbb{R}^2$ o in $\mathbb{C}^2$ a seconda del caso) è che siano linearmente indipendenti, e questa proprietà viene condivisa dalla loro immagine attraverso Q :

$$\alpha Q\mathbf{u}_1 + \beta Q\mathbf{u}_2 = 0 \implies Q(\alpha\mathbf{u}_1 + \beta\mathbf{u}_2) = 0 \implies \alpha\mathbf{u}_1 + \beta\mathbf{u}_2 = 0 \implies \alpha = \beta = 0$$

se $\mathbf{u}_1$ e $\mathbf{u}_2$ sono linearmente indipendenti.

2. Consideriamo un endomorfismo rappresentato dalla matrice A . Si può dimostrare che tutte e sole le matrici che rappresentano lo stesso endomorfismo rispetto alle varie basi di $\mathbb{C}^2$ hanno la forma $B = Q^{-1}AQ$. Q è la matrice detta di *cambiamento di base*, di cui ora richiamiamo la costruzione. Siano $E := (\mathbf{e}_1, \mathbf{e}_2)$, $F := (\mathbf{u}_1, \mathbf{u}_2)$ due basi, e

siano A, B le matrici che rappresentano il medesimo endomorfismo nella base E e F , rispettivamente. Allora $B = Q^{-1}AQ$ dove:

$$\begin{pmatrix} u_{11} & u_{21} \\ u_{12} & u_{22} \end{pmatrix} \quad (2.24)$$

Qui u_{ij} denota la componente di $\mathbf{u}_i$ rispetto a $\mathbf{e}_j$. Più precisamente: $\mathbf{u}_i = u_{i1}\mathbf{e}_1 + u_{i2}\mathbf{e}_2$.

3. Supponiamo che gli autovettori $\mathbf{u}_1$ e $\mathbf{u}_2$ di cui sopra siano ortogonali: $\langle \mathbf{u}_1, \mathbf{u}_2 \rangle = 0$. Allora il cambiamento di base definito da Q conserva i prodotti scalari e quindi in particolare l'ortogonalità: $\langle Q\mathbf{u}_1, Q\mathbf{u}_2 \rangle = 0$. Per dimostrare questo fatto importante, notiamo che in questo caso la matrice Q è *ortogonale*, cioè l'inversa Q^{-1} coincide con la trasposta Q^T . Infatti:

$$\begin{aligned} Q &= \begin{pmatrix} u_{11} & u_{21} \\ u_{12} & u_{22} \end{pmatrix}; & Q^T &= \begin{pmatrix} u_{11} & u_{12} \\ u_{21} & u_{22} \end{pmatrix}; \\ Q^T Q &= \begin{pmatrix} u_{11}^2 + u_{12}^2 & 0 \\ 0 & u_{21}^2 + u_{22}^2 \end{pmatrix} = \begin{pmatrix} \|\mathbf{u}_1\|^2 & 0 \\ 0 & \|\mathbf{u}_2\|^2 \end{pmatrix} \end{aligned}$$

perché la condizione di ortogonalità scritta per componenti è $u_{11}u_{21} + u_{12}u_{22} = 0$. Inoltre possiamo sempre assumere $\|\mathbf{u}_1\|^2 = \|\mathbf{u}_2\|^2 = 1$; altrimenti si considerano i vettori $\mathbf{v}_1 := \mathbf{u}_1/\|\mathbf{u}_1\|$ e $\mathbf{v}_2 := \mathbf{u}_2/\|\mathbf{u}_2\|$ che sono sempre ortogonali e hanno norma 1. Dunque $Q^T Q = I$ da cui segue $Q Q^T = I$ e $Q^T = Q^{-1}$. Pertanto, $\forall \mathbf{u}, \mathbf{v} \in \mathbb{C}^2$:

$$\langle Q\mathbf{u}, Q\mathbf{v} \rangle = \langle Q^T Q\mathbf{u}, \mathbf{v} \rangle = \langle Q^{-1} Q\mathbf{u}, \mathbf{v} \rangle = \langle \mathbf{u}, \mathbf{v} \rangle$$

e ciò dimostra l'asserzione (si ricordi che $\langle \mathbf{u}, Q\mathbf{v} \rangle = \langle Q^T \mathbf{u}, \mathbf{v} \rangle$).

In particolare $\langle Q\mathbf{u}_1, Q\mathbf{u}_2 \rangle = \langle \mathbf{u}_1, \mathbf{u}_2 \rangle$ e l'asserzione è così provata.

4. Il determinante $\det Q$ di una matrice ortogonale Q vale ± 1 . Infatti:

$$1 = \det Q Q^{-1} = \det Q \det Q^T = [\det Q]^2$$

perché l'operazione di trasposizione (scambio delle righe con le colonne) lascia notoriamente invariato il valore del determinante di una matrice quadrata.

Un caso molto semplice in cui si può stabilire subito se esistono due autovettori ortogonali è quello in cui la matrice L è *simmetrica*, cioè coincide con la propria trasposta: $L = L^T$. Ora se $L = \begin{pmatrix} l_{11} & l_{12} \\ l_{21} & l_{22} \end{pmatrix}$ la matrice trasposta L^T , che si ottiene scambiando le righe con le colonne, è $L^T = \begin{pmatrix} l_{11} & l_{21} \\ l_{12} & l_{22} \end{pmatrix}$; Pertanto L è simmetrica se e solo se $l_{12} = l_{21}$. Si ha allora

Proposizione 2.15 *Sia L simmetrica. Allora:*

1. *Gli autovalori λ_1 e λ_2 di L sono reali;*
2. *L ammette sempre due autovettori ortogonali;*
3. *La matrice L è simile ad una matrice diagonale, cioè esiste una matrice ortogonale Q tale che $Q^T L Q = D$, dove $D := \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix}$.*

Dimostrazione.

1. L'equazione caratteristica

$$\det(L - \lambda I) = \lambda^2 - (l_{11} + l_{22})\lambda + l_{11}l_{22} - l_{12}^2$$

ammette le soluzioni

$$\lambda_{1,2} = \frac{1}{2}[l_{11} + l_{22} \pm \sqrt{(l_{11} + l_{22})^2 - 4l_{11}l_{22} + 4l_{12}^2}] = \frac{1}{2}[l_{11} + l_{22} \pm \sqrt{(l_{11} - l_{22})^2 + 4l_{12}^2}] \quad (2.25)$$

che sono entrambe reali.

Ricordiamo poi che essendo $\text{Tr}L := l_{11} + l_{22}$ la *traccia* della matrice L , l'equazione caratteristica può essere abbreviata così:

$$\lambda^2 - \text{Tr}L\lambda + \det L = 0$$

2. Si noti anzitutto che la proprietà di simmetria $L = L^T$ implica $\langle L\mathbf{x}, \mathbf{y} \rangle = \langle \mathbf{x}, L\mathbf{y} \rangle$ per ogni coppia di vettori $\mathbf{x}, \mathbf{y}$, perché $\langle \mathbf{x}, L\mathbf{y} \rangle = \langle L^T \mathbf{x}, \mathbf{y} \rangle = \langle L\mathbf{x}, \mathbf{y} \rangle$. Sia per cominciare $\lambda_1 \neq \lambda_2$. Prendiamo due autovettori $\mathbf{u}_1, \mathbf{u}_2$ corrispondenti agli autovalori λ_1 e λ_2 , rispettivamente. Allora, usando la linearità di L e le relazioni $L\mathbf{u}_1 = \lambda_1\mathbf{u}_1, L\mathbf{u}_2 = \lambda_2\mathbf{u}_2$:

$$\lambda_1 \langle \mathbf{u}_1, \mathbf{u}_2 \rangle = \langle L\mathbf{u}_1, \mathbf{u}_2 \rangle = \langle \mathbf{u}_1, L\mathbf{u}_2 \rangle = \lambda_2 \langle \mathbf{u}_1, \mathbf{u}_2 \rangle$$

Sottraendo membro a membro troviamo $(\lambda_1 - \lambda_2)\langle \mathbf{u}_1, \mathbf{u}_2 \rangle = 0$ da cui $\langle \mathbf{u}_1, \mathbf{u}_2 \rangle = 0$ perché $\lambda_1 \neq \lambda_2$. Dunque autovettori corrispondenti ad autovalori differenti sono ortogonali se L è simmetrica.

3. Siano ora $\mathbf{u}_1 = (u_{11}, u_{12})$, $\mathbf{u}_2 = (u_{21}, u_{22})$ autovettori di norma 1 corrispondenti a λ_1, λ_2 rispettivamente. Abbiamo già visto che ciò equivale alle relazioni seguenti:

$$\begin{cases} u_{11}u_{21} + u_{12}u_{22} = 0 \\ u_{11}^2 + u_{12}^2 = 1 \\ u_{21}^2 + u_{22}^2 = 1 \end{cases} \quad (2.26)$$

Sia ancora $Q = \begin{pmatrix} u_{11} & u_{21} \\ u_{12} & u_{22} \end{pmatrix}$. Per semplificare la notazione poniamo $l_{11} = a$, $l_{12} = l_{21} = b$, $l_{22} = c$. Le condizioni di autovettori $Lu_1 = \lambda_1 u_1$ e $Lu_2 = \lambda_2 u_2$ diventano allora, scritte per esteso:

$$\begin{pmatrix} a & b \\ b & c \end{pmatrix} \begin{pmatrix} u_{11} \\ u_{12} \end{pmatrix} = \lambda_1 \begin{pmatrix} u_{11} \\ u_{12} \end{pmatrix}, \quad \begin{pmatrix} a & b \\ b & c \end{pmatrix} \begin{pmatrix} u_{21} \\ u_{22} \end{pmatrix} = \lambda_2 \begin{pmatrix} u_{21} \\ u_{22} \end{pmatrix},$$

da cui si ottiene, eseguendo i prodotti riga per colonna:

$$\begin{cases} au_{11} + bu_{12} = \lambda_1 u_{11} \\ bu_{11} + cu_{12} = \lambda_1 u_{12} \\ au_{21} + bu_{22} = \lambda_2 u_{21} \\ bu_{21} + cu_{22} = \lambda_2 u_{22} \end{cases} \quad (2.27)$$

Sia ora $D := Q^{-1}TLQ$. Sappiamo che $Q^{-1} = Q^T$, e quindi, usando la (2.26) e la (2.27):

$$\begin{aligned} D &= Q^{-1}TLQ = \begin{pmatrix} u_{11} & u_{12} \\ u_{21} & u_{22} \end{pmatrix} \begin{pmatrix} a & b \\ b & c \end{pmatrix} \begin{pmatrix} u_{11} & u_{21} \\ u_{12} & u_{22} \end{pmatrix} = \\ &= \begin{pmatrix} u_{11} & u_{12} \\ u_{21} & u_{22} \end{pmatrix} \begin{pmatrix} au_{11} + bu_{12} & au_{21} + bu_{22} \\ bu_{11} + cu_{12} & bu_{21} + cu_{22} \end{pmatrix} = \\ &= \begin{pmatrix} u_{11} & u_{12} \\ u_{21} & u_{22} \end{pmatrix} \begin{pmatrix} \lambda_1 u_{11} & \lambda_2 u_{21} \\ \lambda_1 u_{12} & \lambda_2 u_{22} \end{pmatrix} = \\ &= \begin{pmatrix} \lambda_1(u_{11}^2 + u_{12}^2) & \lambda_2(u_{11}u_{21} + u_{12}u_{22}) \\ \lambda_1(u_{11}u_{21} + u_{12}u_{22}) & \lambda_2(u_{21}^2 + u_{22}^2) \end{pmatrix} = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} \end{aligned}$$

e ciò conclude la prova del punto 3 quando $\lambda_1 \neq \lambda_2$.

Se $\lambda_1 = \lambda_2 = \lambda$, per la (2.25) deve necessariamente essere $l_{11} - l_{22} = 0$ e $l_{12} = 0$, da cui $l_{11} = l_{22}$ e $l_{12} = l_{21} = 0$. Pertanto L sarà della forma

$$L = \begin{pmatrix} l & 0 \\ 0 & l \end{pmatrix} = l \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \lambda I$$

Dunque L è diagonale; inoltre dalla (2.25) si deduce $\lambda = \frac{1}{2}(l_{11} + l_{22}) = l$, ossia

$$L = \begin{pmatrix} \lambda & 0 \\ 0 & \lambda \end{pmatrix} = \lambda I$$

Qualsiasi vettore di C^2 è evidentemente autovettore della matrice identità corrispondente all'autovalore 1. Poiché L è un multiplo dell'identità di fattore λ , qualsiasi vettore di C^2 è ugualmente autovettore corrispondente all'autovalore λ . Quindi possiamo prendere, ad esempio, la base canonica e_1 e e_2 come base di autovettori linearmente indipendenti. Ciò prova il punto 3 nel caso $\lambda_1 = \lambda_2 = \lambda$ e si conclude così la dimostrazione della proposizione.

Esercizio 2.40

Data la matrice

$$L := \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$$

trovarne gli autovalori, gli autovettori e la trasformazione ortogonale Q che la riduce a forma diagonale *Soluzione*. La matrice è simmetrica. L'equazione caratteristica è

$$\lambda^2 - 3\lambda + 1 = 0 \implies \lambda_{1,2} = \frac{1}{2}[3 \pm \sqrt{5}]$$

Pertanto la forma diagonale di L è

$$D = \begin{pmatrix} \frac{1}{2}[3 + \sqrt{5}] & 0 \\ 0 & \frac{1}{2}[3 - \sqrt{5}] \end{pmatrix}$$

Per trovare la matrice Q che stabilisce la similitudine $Q^{-1}LQ = D$, che come sappiamo in questo caso sarà ortogonale, basta trovare gli autovettori e metterli in colonna. Le componenti u_{11}, u_{12} dell'autovettore $\mathbf{u}_1$, e quelle u_{21}, u_{22} dell'autovettore $\mathbf{u}_2$ sono le soluzioni dei sistemi lineari omogenei

$$\begin{cases} (2 - \lambda_1)u_{11} + u_{12} = 0 \\ u_{11} + (1 - \lambda_1)u_{12} = 0 \end{cases} \quad \begin{cases} (2 - \lambda_2)u_{21} + u_{22} = 0 \\ u_{21} + (1 - \lambda_2)u_{22} = 0 \end{cases}$$

Per risolverli, fissiamo al solito la prima componente: $u_{11} = u_{21} = 1$. Possiamo farlo perché queste componenti sono necessariamente diverse da 0. Altrimenti, dato che $\lambda_1 \neq 0$, $\lambda_2 \neq 0$, lo sarebbero anche le seconde componenti e quindi $\mathbf{u}_1$ e $\mathbf{u}_2$ non sarebbero autovettori. Allora: $u_{12} = \lambda_1 - 2$, $u_{22} = \lambda_2 - 2$. Quindi gli autovettori sono $\mathbf{u}_1 = \begin{pmatrix} 1 \\ \lambda_1 - 2 \end{pmatrix}$ e $\mathbf{u}_2 = \begin{pmatrix} 1 \\ \lambda_2 - 2 \end{pmatrix}$. Si noti che

$$\langle \mathbf{u}_1, \mathbf{u}_2 \rangle = 1 + (\lambda_1 - 2)(\lambda_2 - 2) = 1 + \frac{1}{4}(\sqrt{5} - 1)(-1 - \sqrt{5}) = 0$$

come dovevamo aspettarci dato che i vettori $\mathbf{u}_1$ e $\mathbf{u}_2$ devono essere l'ortogonali. Quindi, per definizione:

$$Q = \begin{pmatrix} 1 & 1 \\ \frac{1}{2}(\sqrt{5} - 1) & \frac{1}{2}(-1 - \sqrt{5}) \end{pmatrix}$$

3 Dinamica discreta ed iterazione di matrici 2×2

Sia data una trasformazione lineare L da $\mathbb{R}^2$ in sè, che identificheremo al solito con la sua matrice 2×2 sulla base canonica.

Definizione 3.35 *Dicesi dinamica in $\mathbb{R}^2$, o piana, l'applicazione*

$$\mathbf{x} \mapsto L^n \mathbf{x}, \quad n = 0, 1, 2, \dots$$

che ad ogni vettore $\mathbf{x} \in \mathbb{R}^2$ (detto dato iniziale) fa corrispondere la successione $L^n \mathbf{x}$ di punti in $\mathbb{R}^2$. Il sottoinsieme di $\mathbb{R}^2$ definito da $\bigcup_{n \geq 0} L^n \mathbf{x}$ si dice orbita del punto $\mathbf{x}$.

Osservazione 3.54

1. Se la matrice L è invertibile possiamo definire la matrice L^{-n} e quindi la dinamica per ogni $n \in \mathbb{Z}$. In tal caso distingueremo al solito fra il futuro ($n > 0$) e il passato ($n < 0$).
2. Denotiamo $\mathbf{y}(n) = \begin{pmatrix} y_1(n) \\ y_2(n) \end{pmatrix}$ il punto dell'orbita all'istante n , cioè:

$$\begin{pmatrix} y_1(n) \\ y_2(n) \end{pmatrix} := \begin{pmatrix} (L^n \mathbf{x})_1 \\ (L^n \mathbf{x})_2 \end{pmatrix}$$

(omettiamo per semplicità di denotare anche la dipendenza dal punto iniziale $\mathbf{x}$). Eliminando n fra le due equazioni si ottiene una relazione fra y_1 e y_2 . Questa relazione definisce una curva in $\mathbb{R}^2$ e pertanto rappresenta l'*equazione cartesiana* dell'orbita. Equivalentemente, si tratta del luogo dei punti successivamente descritti dall'evoluzione discreta di ogni fissato dato iniziale.

Abbiamo visto che il caso di gran lunga più semplice di dinamica discreta è rappresentato dalla relazione di ricorrenza a due termini ad incremento costante $a_{n+1} = \lambda a_n$, con la condizione iniziale $a_0 = x \in \mathbb{R}$. Il solo punto fisso è $x = 0$, che è stabile ed esponenzialmente attrattivo per $|\lambda| < 1$ mentre è instabile ed esponenzialmente repulsivo per $|\lambda| > 1$. La convergenza a 0 è monotona per $0 < \lambda < 1$, mentre è oscillatoria con ampiezza delle oscillazioni esponenzialmente decrescente per $-1 < \lambda < 0$. Allo stesso modo, la divergenza è monotona per $1 < \lambda$, mentre è oscillatoria con ampiezza delle oscillazioni esponenzialmente crescente per $\lambda < -1$. Vogliamo ora fare vedere che la riduzione delle matrici 2×2 a forma canonica compiuta nel paragrafo precedente permette di fare rientrare in questo caso anche la dinamica bidimensionale sopra definita. Cominciamo al solito con degli esempi.

Esempio 3.39

1. Sia $L_1 := \begin{pmatrix} 2 & 0 \\ 0 & \frac{1}{2} \end{pmatrix}$. L_1 è diagonale e quindi l'orbita di qualsiasi dato iniziale $\mathbf{x} = (x_1, x_2)$ si calcola immediatamente:

$$\mathbf{y}(n) = L_1^n \mathbf{x} = \begin{pmatrix} 2^n & 0 \\ 0 & 2^{-n} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 2^n x_1 \\ 2^{-n} x_2 \end{pmatrix}$$

cioè $y_1(n) = 2^n x_1$, $y_2(n) = 2^{-n} x_2$. Poiché la matrice L_1 è diagonale, non c'è alcun accoppiamento fra le variabili x_1 e x_2 e infatti la dinamica discreta lungo una delle due direzioni coordinate è completamente indipendente dall'altra. Lungo l'asse x_1 si ha repulsione con velocità esponenziale nel futuro per $x_1 \neq 0$ e lungo la direzione x_2 attrazione con velocità esponenziale per $x_2 \neq 0$. L'origine $(0, 0)$ è un punto fisso. Il comportamento sui due assi ortogonali si inverte scambiando il futuro ($n > 0$) con il passato ($n < 0$). Possiamo ora eliminare n nelle due equazioni $y_1 = 2^n x_1$, $y_2 = 2^{-n} x_2$ per trovare

$$y_1 y_2 = x_1 x_2$$

Questa è l'equazione cartesiana di una famiglia di iperboli equilateri (ce n'è una per ogni scelta del prodotto $x_1 x_2$ dei dati iniziali) sempreché $x_1 \neq 0$ e $x_2 \neq 0$. Dunque le orbite sono tutte iperboli equilateri. Per $x_1 = 0$ l'orbita è l'asse x_2 , lungo il quale ogni condizione iniziale converge esponenzialmente all'origine, e per $x_2 = 0$ l'orbita è l'asse x_1 , lungo il quale ogni condizione iniziale non nulla si allontana esponenzialmente dall'origine.

2. Sia ora $L_2 := \begin{pmatrix} \frac{1}{2} & 0 \\ 0 & \frac{1}{3} \end{pmatrix}$. Anche qui le variabili x_1 e x_2 sono indipendenti perché la matrice è diagonale e si calcola subito

$$\mathbf{y}(n) = L_2^n \mathbf{x} = \begin{pmatrix} 2^{-n} & 0 \\ 0 & 3^{-n} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 2^{-n} x_1 \\ 3^{-n} x_2 \end{pmatrix}$$

cioè $y_1(n) = 2^{-n} x_1$, $y_2(n) = 3^{-n} x_2$. Anche qui l'origine è un punto fisso, e si ha attrazione verso l'origine con velocità esponenziale (nel futuro) lungo entrambe le direzioni. Troviamo ora l'equazione cartesiana dell'orbita. Dalla $y_1 = 2^{-n} x_1$ deduciamo $-n = \ln \frac{y_1}{x_1} / \ln 2$ e sostituendo questo valore di $-n$ nella $y_2 = 3^{-n} x_2$ troviamo

$$y_2 = 3^{\ln \frac{y_1}{x_1} / \ln 2} x_2 = (3^{\ln \frac{y_1}{x_1}})^{1/\ln 2} x_2 = ((e^{\ln 3})^{\ln \frac{y_1}{x_1}})^{1/\ln 2} x_2 = (e^{\ln \frac{y_1}{x_1}})^{\ln 3 / \ln 2} x_2 = \left(\frac{y_1}{x_1} \right)^{\ln 3 / \ln 2} x_2$$

da cui infine

$$y_2 = C(y_1)^{\ln 3 / \ln 2}, \quad C(x_1, x_2) = (x_1)^{-\ln 3 / \ln 2} x_2$$

Poiché $\ln 3 / \ln 2 > 1$, si tratta di una curva passante per l'origine, percorsa verso l'origine medesima per quanto visto sopra.

3. Sia ora $L_3 := \begin{pmatrix} 0 & -\frac{1}{2} \\ \frac{1}{2} & 0 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$. Stavolta le variabili non sono più indipendenti. Il calcolo dell'iterazione L_3^n l'abbiamo già fatto nel punto 3 dell'Esercizio (2.38). La (2.8) per $a = 1/2$ dà immediatamente

$$\mathbf{y}_{2p} = L_3^{2p} \mathbf{x} = 2^{-2p} \mathbf{x}; \quad \mathbf{y}_{2p+1} = 2^{-2p-1} J \mathbf{x}$$

da cui

$$\begin{cases} y_1(2p) = 2^{-2p} x_1 \\ y_2(2p) = 2^{-2p} x_2 \end{cases} \quad \begin{cases} y_1(2p+1) = -2^{-2p-1} x_2 \\ y_2(2p+1) = 2^{-2p-1} x_1 \end{cases}$$

Queste equazioni mostrano che ogni condizione iniziale esegue a intervalli di tempo alterni una conversione verso l'origine di passo $1/2$ lungo la bisettrice $y_1 = y_2$ seguita da una rotazione antioraria di ampiezza $\pi/2$ che scambia le ascisse con le ordinate. Dunque ogni condizione iniziale diversa dal punto fisso all'origine converge verso l'origine medesima lungo una spirale logaritmica di passo $1/2$.

4. Sia ora $L_4 := \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$. Qui notiamo che:

$$\mathbf{y}_n = L_4 \mathbf{y}_{n-1} = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} y_1(n-1) \\ y_2(n-1) \end{pmatrix} = \begin{pmatrix} 2y_1(n-1) + y_2(n-1) \\ y_1(n-1) + y_2(n-1) \end{pmatrix}$$

definiamo ora la successione a_n le modo seguente: $a_n := y_2(n)$; $a_{n-1} = y_1(n)$. Allora la formula precedente porge

$$\begin{pmatrix} a_{n+1} \\ a_n \end{pmatrix} = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} a_{n-1} \\ a_{n-2} \end{pmatrix} \quad \text{ovvero} \quad \begin{cases} a_{n+1} = 2a_{n-1} + a_{n-2} \\ a_n = a_{n-1} + a_{n-2} \end{cases}$$

da cui, poiché $a_{n-2} = a_n - a_{n-1}$, possiamo concludere che

$$\begin{cases} a_{n+1} = a_n + a_{n-1} \\ a_n = a_{n-1} + a_{n-2} \end{cases}$$

In altre parole, la successione a_n soddisfa la ricorrenza di Fibonacci. Ritroviamo ora per questa via l'espressione esplicita (4.15) della successione di Fibonacci. Anzitutto

occorrerà imporre le condizioni iniziali $a_0 = 0, a_1 = 1$. tradotto nel linguaggio vettoriale che stiamo usando, questo equivale a considerare il dato iniziale $\mathbf{y}_0 \equiv \mathbf{x} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$. Sia $\mathbf{u}_1$ autovettore di L_4 corrispondente all'autovalore λ_1 . Allora $L_4^n \mathbf{u}_1 = \lambda_1^n \mathbf{u}_1$. Quindi basterà: sviluppare $\mathbf{x}$ sulla base ortogonale degli autovettori $\mathbf{u}_1$ e $\mathbf{u}_2$ di L_4 trovata nell'esercizio 2.40. Poniamo $\mathbf{x} = \alpha_1 \mathbf{u}_1 + \alpha_2 \mathbf{u}_2$. Allora le costanti α_1 e α_2 saranno le soluzioni uniche del sistema lineare non omogeneo

$$\begin{cases} \alpha_1 + \alpha_2 = 1 \\ (\lambda_1 - 2)\alpha_1 + (\lambda_2 - 1)\alpha_2 = 0 \end{cases}$$

Pertanto $\alpha_1 = \frac{1}{\lambda_1 + \lambda_2 - 1} = \frac{1}{2}$, $\alpha_2 = \frac{1}{2}$ e:

$$\begin{aligned} L_4^n \mathbf{x} &= \alpha_1 (\lambda_1^n \mathbf{u}_1 + \lambda_2^n \mathbf{u}_2) = \frac{1}{2} \begin{pmatrix} \lambda_1^n + \lambda_2^n \\ \lambda_1^n (\lambda_1 - 2) + \lambda_2^n (\lambda_2 - 1) \end{pmatrix} \\ &= \frac{1}{2^{n+1}} \begin{pmatrix} (\sqrt{5} + 3)^n + (\sqrt{5} - 3)^n \\ \frac{1}{2}(\sqrt{5} - 1)[(\sqrt{5} + 3)^n - (\sqrt{5} - 3)^n] \end{pmatrix} \end{aligned}$$

Scrivendo $\frac{\sqrt{5} + 3}{2} = \frac{\sqrt{5} + 1}{2} + 1$, $\frac{\sqrt{5} - 3}{2} = \frac{\sqrt{5} - 1}{2} - 1$ è ora facile vedere che la successione $a_n = \frac{1}{2^{n+1}} \frac{1}{2} (\sqrt{5} - 1)[(\sqrt{5} + 3)^n - (\sqrt{5} - 3)^n]$ coincide con la (4.15).

Come si è accennato in precedenza questi esempi, che permettono di determinare immediatamente le orbite e il loro andamento, sono solo apparentemente particolari. In realtà, se si esclude per semplicità il caso triangolare della forma canonica matrice L , la soluzione esplicita dell'iterazione della matrice 2×2 può essere sempre riportata a questi casi. Si noti che ovviamente l'origine $\mathbf{x} = 0$ è comunque il solo punto fisso per ogni trasformazione L .

Proposizione 3.16 *Supponiamo che la trasformazione lineare L da $\mathbb{R}^2$ in sé ammetta due autovettori linearmente indipendenti $\mathbf{u}_1, \mathbf{u}_2$, in corrispondenza agli autovalori λ_1 e λ_2 reali (non necessariamente distinti) o complessi coniugati. Sia $\mathbf{x} \in \mathbb{R}^2$ un vettore non nullo arbitrario, e sia $\mathbf{y}(n) := L^n \mathbf{x}$. Allora:*

1. Se gli autovalori λ_1 e λ_2 sono reali si ha

$$\mathbf{y}(n) = L^n \mathbf{x} = \alpha_1(\mathbf{x}) \lambda_1^n \mathbf{u}_1 + \alpha_2(\mathbf{x}) \lambda_2^n \mathbf{u}_2 \quad (3.28)$$

dove le costanti $\alpha_1(\mathbf{x})$ e $\alpha_2(\mathbf{x})$ sono le componenti del dato iniziale $\mathbf{x}$ lungo la base degli autovettori, e cioè le soluzioni del sistema lineare non omogeneo

$$\begin{cases} \alpha_1 u_{11} + \alpha_2 u_{21} = x_1 \\ \alpha_1 u_{12} + \alpha_2 u_{22} = x_2 \end{cases} \quad \text{ovvero} \quad Q \begin{pmatrix} \alpha_1 \\ \alpha_2 \end{pmatrix} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \quad (3.29)$$

2. Se gli autovalori λ_1 e λ_2 sono complessi coniugati, $\lambda_{1,2} := \lambda \pm i\mu$ si ha

$$\mathbf{y}(n) = L^n \mathbf{x} = \rho^n (\alpha_1(\mathbf{x}) e^{in\phi} \mathbf{u}_1 + \alpha_2(\mathbf{x}) e^{-in\phi} \mathbf{u}_2) \quad (3.30)$$

Qui le costanti $\alpha_1(\mathbf{x})$ e $\alpha_2(\mathbf{x})$ sono ancora le componenti del dato iniziale $\mathbf{x}$ lungo la base degli autovettori, e cioè le soluzioni del sistema (3.29), e ρ , ϕ sono modulo e fase del numero complesso $\lambda + i\mu$: $\rho = \sqrt{\lambda^2 + \mu^2}$, $\phi = \arctg \frac{\mu}{\lambda}$.

Dimostrazione. per ipotesi gli autovettori $\mathbf{u}_1$, $\mathbf{u}_2$, essendo linearmente indipendenti, formano una base in $\mathbb{R}^2$. Quindi ogni vettore $\mathbf{x} \in \mathbb{R}^2$ può essere scritto sotto la forma $\mathbf{x} = \alpha_1(\mathbf{x})\mathbf{u}_1 + \alpha_2(\mathbf{x})\mathbf{u}_2$ dove α_1 e α_2 sono le soluzioni del sistema lineare (3.29). Poiché $\mathbf{u}_1$, $\mathbf{u}_2$ sono autovettori di L corrispondenti agli autovalori λ_1 e λ_2 facendo agire n volte L su $\mathbf{x}$ si ottiene la formula (3.28). Lo stesso ragionamento conduce alla formula (3.30) osservando che in questo caso $\lambda_1^n = (\lambda + i\mu)^n = \rho^n e^{in\phi}$, $\lambda_2^n = \overline{\lambda_1}^n = \rho^n e^{-in\phi}$. Ciò conclude la prova della proposizione.

Da questa rappresentazione esplicita possiamo dedurre immediatamente tutte le proprietà della dinamica discreta $\mathbf{x} \mapsto L^n \mathbf{x}$, comprese quelle di stabilità, convergenza o divergenza esponenziale, descrizione qualitativa delle orbite, ecc. Si noti anzitutto che le definizioni di punto fisso, stabilità, attrattività, repulsività ecc. date nella Definizione 5.18 si applicano in questo caso senza alcuna variazione.

Corollario 3.1 *Nelle ipotesi della proposizione precedente:*

Nel caso 1 (autovalori di L reali) il punto fisso $\mathbf{x} = 0$ è

1. *Stabile se $|\lambda_1| \leq 1$ e $|\lambda_2| \leq 1$;*
2. *Globalmente attrattivo se $|\lambda_1| < 1$ e $|\lambda_2| < 1$;*
3. *Instabile se $|\lambda_1| > 1$ oppure $|\lambda_2| > 1$; globalmente repulsivo se $|\lambda_1| > 1$ e $|\lambda_2| > 1$.*

Nel caso 2 (autovalori di L complessi coniugati) il punto fisso $\mathbf{x} = 0$ è

1. *Stabile se $\rho \leq 1$;*
2. *Globalmente attrattivo se $|\rho| < 1$;*
3. *Globalmente repulsivo se $\rho > 1$ e $|\lambda_2| > 1$.*

Dimostrazione.

Caso 1. L'espressione esplicita della soluzione data dalla (3.28) permette di dedurre immediatamente le affermazioni. Infatti nel sottocaso 1, siccome ovviamente $|\alpha_k(\mathbf{x})| \leq \|\mathbf{x}\|$, dato $\epsilon > 0$ esiste sicuramente $\delta(\epsilon) > 0$ tale che $\|\mathbf{y}(n)\| < \epsilon$ se $\|\mathbf{x}\| < \delta(\epsilon)$. L'affermazione del sottocaso 2 è evidente perché entrambe le successioni λ_1^n e λ_2^n convergono a 0 con velocità esponenziale; il sottocaso 3 anche perché almeno una, o entrambe, queste successioni divergono esponenzialmente. Allo stesso modo si verifica il caso 2, e ciò conclude la dimostrazione del corollario.

Osservazione 3.55

1. Consideriamo il caso 1, quello degli autovalori reali, nel sottocaso instabile, ma non globalmente repulsivo. Allora uno degli autovalori è in valore assoluto maggiore di 1 e l'altro minore di 1. Ad esempio, sia $|\lambda_1| > 1$ e $|\lambda_2| < 1$. Allora la (3.28) mostra immediatamente che se $\alpha_2(\mathbf{x}) = 0$, cioè il dato iniziale ha componente nulla sul sottospazio sotteso dall'autovettore $\mathbf{u}_2$, il dato iniziale $\mathbf{x}$ viene in realtà respinto all'infinito con velocità esponenziale, mentre se, viceversa, ha componente nulla sul sottospazio sotteso dall'autovettore $\mathbf{u}_1$ esso viene attratto dal punto fisso all'origine con velocità esponenziale. Per la (3.29) possiamo caratterizzare facilmente questi insiemi di dati iniziali: si ha infatti

$$\begin{pmatrix} \alpha_1 \\ \alpha_2 \end{pmatrix} = Q^{-1} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$$

Esplicitamente, risolvendo il sistema (3.29) con la regola di Cramer:

$$\alpha_1 = \frac{u_{11}u_{22}x_1 - u_{21}u_{12}x_2}{u_{11}u_{22} - u_{21}u_{12}}; \quad \alpha_2 = \frac{u_{11}u_{22}x_2 - u_{21}u_{12}x_1}{u_{11}u_{22} - u_{21}u_{12}}$$

Dunque si ha attrattività esponenziale dell'origine per tutti i dati iniziali $\mathbf{x}$ tali che $u_{22}x_1 - u_{21}u_{12}x_2$, mentre si ha repulsività esponenziale se $\mathbf{x}$ è tale che $u_{11}u_{22}x_2 - u_{21}u_{12}x_1$. Si tratta come si vede subito di due rette ortogonali fra di loro. Quindi,

secondo la definizione 5.11, chiameremo queste rette *varietà* o *insieme stabile*, denotato W_s , o *instabile*, denotato W_u , rispettivamente.

Chiudiamo questo mostrando che l'equivalenza fra le relazioni di ricorrenza a tre termini e la dinamica definita dall'iterazione di una matrice 2×2 , già constatata nel caso particolare dei numeri di Fibonacci, vale nel caso della ricorrenza a tre termini generale (a coefficienti costanti) $a_{n+2} = Aa_{n+1} + Ba_n$. Si ha

Proposizione 3.17 *La dinamica discreta generata dalla ricorrenza a tre termini*

$$a_{n+2} = Aa_{n+1} + Ba_n, \quad A \in \mathbb{R}, B \in \mathbb{R}$$

coincide con la dinamica discreta generata dall'iterazione della matrice:

$$A := \begin{pmatrix} A(1+B) & B^2 \\ A & B \end{pmatrix}$$

nel senso che se $a_n(p, q)$ è la soluzione unica della ricorrenza con le condizioni iniziali $a_0 = p$, $a_1 = q$ allora essa coincide con la successione di vettori $b_n(p, q)$ definita nel modo seguente

$$\begin{pmatrix} b_{n+1}(p, q) \\ b_n(p, q) \end{pmatrix} := A^n \begin{pmatrix} q \\ p \end{pmatrix}$$

Dimostrazione. Basta dimostrare che $A^n \begin{pmatrix} q \\ p \end{pmatrix}$ genera la ricorrenza $a_{n+2} = Aa_{n+1} + Ba_n$. L'identità delle condizioni iniziali garantisce poi l'identità delle soluzioni. Sia:

$$\begin{pmatrix} b_{n+1} \\ b_n \end{pmatrix} = \begin{pmatrix} A(1+B) & B^2 \\ A & B \end{pmatrix} \begin{pmatrix} b_{n-1} \\ b_{n-2} \end{pmatrix}$$

da cui

$$\begin{aligned} b_{n+1} &= A(1+B)b_n + B^2b_{n-2} \\ b_n &= Ab_{n-1} + Bb_{n-2} \end{aligned}$$

Ricavando b_{n-2} dalla seconda e sostituendo nella prima ricaviamo $b_{n+1} = Ab_n + Bb_{n-1}$ e ciò dimostra l'asserzione.

4 Dinamica discreta ed automorfismi del toro

Nel paragrafo precedente abbiamo trattato in maniera esauriente la dinamica discreta bidimensionale nel piano. In sostanza non succede niente di veramente nuovo rispetto al caso

unidimensionale. Si tratta solo di scegliere gli assi giusti per ridurre il caso bidimensionale ad un semplice prodotto cartesiano di due casi unidimensionali. La situazione cambia drasticamente se al posto del piano consideriamo la dinamica discreta le cui orbite giacciono su superfici compatte, come ad esempio il toro bidimensionale $\mathbb{T}^2$.

Sia S^1 la circonferenza di raggio 1. Abbiamo già notato che per conoscere la posizione di ciascun punto sulla circonferenza è sufficiente conoscere la sua coordinata angolare ϕ che prende valori fra 0 e 2π : $0 \leq \phi \leq 2\pi$ (la corrispondenza fra punti della circonferenza e valori dell'angolo è così biunivoca tranne per il punto 0 che viene identificato con 2π). Per costruzione l'angolo ϕ è definito a meno di un multiplo intero di 2π : aggiungere o togliere a un tale angolo non cambia la posizione del punto sulla circonferenza. È quindi naturale porre una definizione puramente algebrica della circonferenza S^1 :

$$S^1 := \mathbb{R} \pmod{\mathbb{Z}}$$

Possiamo allora definire anche il toro bidimensionale $\mathbb{T}^2$ per via puramente algebrica:

Definizione 4.36 *Il toro $\mathbb{T}^2$ è il sottoinsieme di $\mathbb{R}^2$ definito algebricamente come $S^1 \times S^1$.*

In formule:

$$\mathbb{T}^2 = S^1 \times S^1 = \mathbb{R} \pmod{\mathbb{Z}} \times \mathbb{R} \pmod{\mathbb{Z}} = \mathbb{R}^2 \pmod{2\pi\mathbb{Z}^2}$$

Osservazione 4.56

1. I punti del toro $\mathbb{T}^2$ saranno quindi individuati da due coordinate angolari ϕ_1, ϕ_2 : $0 \leq \phi_1 \leq 2\pi, 0 \leq \phi_2 \leq 2\pi$;
2. L'operazione di quozientazione definisce la proiezione naturale π da $\mathbb{R}^2$ su $\mathbb{T}^2$: dato $(x_1, x_2) \in \mathbb{R}^2$ definiamo l'applicazione π da $\mathbb{R}^2$ a $\mathbb{T}^2$ nel modo seguente:

$$\pi(x_1, x_2) = [x_1, x_2] := (x_1 \pmod{N_1}, x_2 \pmod{N_2}) \quad (4.31)$$

dove N_1 e N_2 sono interi tali che $x_1 + N_1 \in [0, 2\pi], x_2 + N_2 \in [0, 2\pi]$. Si noti anche che la distanza euclidea nel piano induce in modo natural una nozione di distanza euclidea sul toro:

$$d(\phi_1, \phi_2) = \sqrt{(\phi_{11} - \phi_{21})^2 + (\phi_{12} - \phi_{22})^2} \quad (4.32)$$

3. La definizione algebrica di S^1 ammette la seguente interpretazione geometrica: la circonferenza è l'intervallo $[0, 2\pi]$ con gli estremi identificati. Allo stesso modo possiamo interpretare il toro $\mathbb{T}^2$ come il quadrato di lato 2π con gli estremi identificati (vedi figura 4.1): $\mathbb{T}^2 = [0, 2\pi]^2$

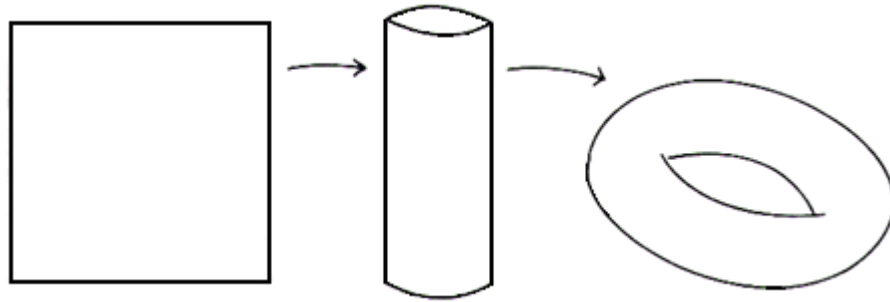


Figure 6.1: Costruzione del toro a partire dal quadrato.

Una differenza molto importante fra $\mathbb{R}^2$ e $\mathbb{T}^2$ è che in $\mathbb{T}^2$ non possiamo più identificare gli endomorfismi, ed in particolare gli automorfismi, con le matrici 2×2 . L'azione della matrice L sul vettore di componenti $\phi = (\phi_1, \phi_2)$ deve conservare la quozientazione se si vuole che L sia ancora in $\mathbb{T}^2$. Pertanto essa non potrà avere coefficienti reali qualunque. Premettiamo allo scopo la definizione di una classe particolare di matrici.

Definizione 4.37 Si denota $SL(2; \mathbb{R}) \subset GL(2; \mathbb{R})$ l'insieme delle matrici 2×2 ad elementi reali unimodulari, cioè con determinante che vale ± 1 . Si denota poi $SL(2; \mathbb{Z}) \subset SL(2; \mathbb{R})$ l'insieme delle matrici 2×2 unimodulari a coefficienti interi.

Esempio 4.40

La matrice $L = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$ che genera la ricorrenza di Fibonacci è in $SL(2; \mathbb{Z})$ perché $\det L = 1$. In generale le matrici $L = \begin{pmatrix} 2g & 4g^2 - 1 \\ 1 & 2g \end{pmatrix}$ appartengono a $SL(2; \mathbb{Z}) \forall g \in \mathbb{Z}$ perché $\det L = 1$.

Esercizio 4.41

Dimostrare che $SL(2; \mathbb{R})$ è un sottogruppo proprio di $GL(2; \mathbb{R})$ e che $SL(2; \mathbb{Z})$ è un sottogruppo proprio di $SL(2; \mathbb{R})$.

Possiamo allora caratterizzare la relazione fra automorfismi lineari del toro e matrici.

Proposizione 4.18 *La matrice 2×2 L genera un automorfismo (cioè una trasformazione lineare invertibile) del toro $\mathbb{T}^2$ in sé se e solo se appartiene a $\mathrm{SL}(2; \mathbb{Z})$.*

Dimostrazione Sia $L \in \mathrm{SL}(2; \mathbb{Z})$, $\phi \in \mathbb{T}^2$. Allora L è chiaramente un automorfismo di $\mathbb{T}^2$: infatti $L\phi \in \mathbb{T}^2$ dato che una matrice a coefficienti interi conserva la quozientazione di $\mathbb{R}^2$ tramite $\mathbb{Z}^2$; inoltre (vedasi l'esercizio precedente) anche $L^{-1} \in \mathrm{SL}(2; \mathbb{Z})$. Sia ora $\mathcal{L}$ un automorfismo lineare di $\mathbb{T}^2$. Vogliamo mostrare che esiste una matrice $L \in \mathrm{SL}(2; \mathbb{Z})$ tale che $\mathcal{L}(\phi) = L\phi$, dove al solito $\phi = \begin{pmatrix} \phi_1 \\ \phi_2 \end{pmatrix}$. Sappiamo, per il medesimo ragionamento della Proposizione 1.1, che la linearità implica l'esistenza di costanti $c_{ik} : i, k = 1, 2$ tali che $\mathcal{L}(\phi)_1 = c_{11}\phi_1 + c_{12}\phi_2$, $\mathcal{L}(\phi)_2 = c_{21}\phi_1 + c_{22}\phi_2$. Affinché $\mathcal{L}(\phi) \in \mathbb{T}^2$ occorre però che i coefficienti c_{ik} siano interi e ciò conclude la prova della proposizione.

Una proprietà molto notevole delle mappe $L \in \mathrm{SL}(2; \mathbb{Z})$ è che la loro azione conserva le aree. Se $V \subset \mathbb{T}^2$ è un qualsiasi rettangolo (aperto, chiuso, semiaperto a destra, ecc;) cioè $V := \{(\phi_1, \phi_2) \in \mathbb{T}^2 \mid a \leq \phi_1 \leq b; c \leq \phi_2 \leq d\}$ dove il simbolo $\leq$ può dovunque essere sostituito da $<$, conveniamo di indicare con $L(V)$ l'immagine di V attraverso L , cioè: $L(V) := \{(\phi_1, \phi_2) \in \mathbb{T}^2 \mid L^{-1}(\phi_1, \phi_2) \in V\}$. Denotiamo poi $\mu(V)$ l'area di V . Ovviamente $\mu(V) = (b-a)(d-c)$. Ora si ha

Proposizione 4.19 *L conserva le aree (a meno del segno). In altre parole l'area $\mu(L(V))$ esiste ed è uguale all'area di V a meno del segno: $\mu(L(V)) = \pm\mu(V)$.*

Dimostrazione. Senza ledere la generalità possiamo assumere $a = c = 0$, di modo che V diventa il rettangolo vertici $P_1 = (0, 0)$, $P_2 = (0, a)$, $P_3 = (a, d)$, $P_4 = (d, 0)$. Allora $L(V)$ sarà il parallelogramma di vertici $LP_1 \pmod{(2\pi)^2}$, $LP_2 \pmod{(2\pi)^2}$, $LP_3 \pmod{(2\pi)^2}$, $LP_4 \pmod{(2\pi)^2}$. D'altra parte $\mu(L(V))$ è l'area del parallelogramma che ha per vertici questi punti; pertanto, dato che il determinante del prodotto vale il prodotto dei determinanti, e che $\mu(V)$ è il volume del rettangolo di vertici $P_1, \dots, P_4$:

$$\mu(L(V)) = \det L \cdot \det \begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix} = \pm\mu(V)$$

e ciò conclude la dimostrazione.

Questo risultato mostra immediatamente come la dinamica discreta su $\mathbb{T}^2$ definita dall'iterazione delle matrici $L \in \mathrm{SL}(2; \mathbb{Z})$ abbia carattere radicalmente diverso da quella definita

su $\mathbb{R}^2$ dall'iterazione di matrici reali. Infatti qui la conservazione delle aree impedisce l'attrattività o la repulsività del punto fisso all'origine. Infatti in questi casi l'area di un qualsiasi rettangolo attorno all'origine dovrebbe o contrarsi o dilatarsi. Un'altra differenza qualitativa molto marcata fra le due dinamiche è messa in luce dal risultato seguente:

Proposizione 4.20 *Sia $L \in \text{SL}(2; \mathbb{Z})$, e si denoti $\text{Per}(L)$ l'insieme dei punti periodici per la la dinamica discreta su $\mathbb{T}^2$ definita dall'iterazione di L . Allora $\text{Per}(L)$ è denso in $\mathbb{T}^2$.*

Dimostrazione. Sia $\phi \in \mathbb{T}^2$ un punto qualsiasi a coordinate razionali. Passando al minimo comune denominatore possiamo sempre scrivere ϕ sotto la forma $\phi = \left(\frac{a}{k}, \frac{b}{k} \right)$, dove a, b e k sono interi con $0 \leq a < k$ e $0 \leq b < k$ (la limitazione a questi valori di a e b è dovuta alla quozientazione: $\frac{k}{k} = 0 \bmod 2\pi$, $\frac{k+1}{k} = \frac{1}{k} \bmod 2\pi, \dots$). Questi punti sono densi in $\mathbb{T}^2$ perché lo sono tutti i punti a coordinate razionali. Dimostriamo adesso che ϕ è periodico con periodo che non supera k^2 . In altre parole esiste $p \leq k^2$ tale che $L^p \phi = \phi$. Per dimostrare questa affermazione cominciamo col notare che i punti di $\mathbb{T}^2$ della forma $\phi = \left(\frac{a}{k}, \frac{b}{k} \right)$ con $0 \leq a < k$ e $0 \leq b < k$ sono esattamente k^2 . Inoltre, poiché gli elementi di L sono interi, anche le componenti di $L\phi$ sono di questa forma. Pertanto l'effetto dell'azione della matrice L è quello di generare una permutazione di questi punti. Esisteranno quindi degli interi i e j con $|i - j| \leq k^2$ tali che $L^i \phi = L^j \phi$. Facendo agire L^{-i} su questa uguaglianza otteniamo $L^{j-i} \phi = \phi$ e quindi ϕ è un punto periodico di periodo che non supera k^2 . Ciò dimostra l'asserzione.

Esempio 4.41

Consideriamo la nostra vecchia amica, la matrice $L = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}$ che genera la ricorrenza di Fibonacci. Per ragioni che diventeranno chiare fra poco, la ribattezziamo *gatto di Arnold*¹. Essa ammette l'origine $(0, 0)$ come il solo punto fisso. Per vedere questo fatto, e trovare dei punti periodici, scriviamo anzitutto le equazioni del punto fisso. Scrivendo per componenti l'equazione $L\phi = \phi \bmod \mathbb{Z}^2$ si ha

$$\begin{cases} 2\phi_1 + \phi_2 = \phi_1 + N_1 \\ \phi_1 + \phi_2 = \phi_2 + N_2 \end{cases}$$

¹Vladimir I. Arnold (Odessa 1937), Professore all'Università di Mosca nonché all'Università di Parigi IX, è uno dei più illustri matematici contemporanei.

dove N_1 e N_2 sono interi. Si vede subito che il solo punto fisso è $(0, 0)$. D'altra parte si vede anche che $L \begin{pmatrix} 1/2 \\ 1/2 \end{pmatrix} = \begin{pmatrix} 1/2 \\ 0 \end{pmatrix}$, $L \begin{pmatrix} 1/2 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 1/2 \end{pmatrix}$, $L \begin{pmatrix} 0 \\ 1/2 \end{pmatrix} = \begin{pmatrix} 1/2 \\ 1/2 \end{pmatrix}$. Dunque $(1/2, 1/2)$ è un punto periodico di periodo 3.

Queste proprietà sono condivise dalla dinamica discreta generata da ogni matrice in $L \in \text{SL}(2; \mathbb{Z})$. Una sottoclasse di queste ha un comportamento caotico come ora vedremo.

5 Dinamica iperbolica sul toro e caos

Riprendiamo anzitutto la nozione di iperbolicità già introdotta per le mappe dell'intervallo e la dinamica bidimensionale in $\mathbb{R}^2$.

Definizione 5.38 *La matrice $L \in \text{SL}(2; \mathbb{Z})$ si dice iperbolica (e l'automorfismo corrispondente di $\mathbb{T}^2$ iperbolico) se i suoi autovalori sono diversi da ± 1 .*

Osservazione 5.57

Poiché $\det L = \pm 1$, deve necessariamente essere $\lambda_1 \lambda_2 \neq 1$, ossia $|\lambda_1| > 1$ e $|\lambda_2| < 1$ o viceversa. Equivalentemente, dato che $\text{Tr } L = \lambda_1 + \lambda_2$, deve essere $|\text{Tr } L| \neq 2$.

Abbiamo visto che il solo punto fisso della dinamica discreta definita dagli automorfismi di $\mathbb{T}^2$, che si trova all'origine, non può essere nè attrattivo nè repulsivo. Ora se l'automorfismo è iperbolico risulta facile caratterizzarne l'insieme stabile e l'insieme instabile. Per quanto abbiamo visto nel caso di un automorfismo iperbolico del piano $\mathbb{R}^2$ in sè sappiamo che gli insiemi stabile ed instabile W_s e W_u sono rette ortogonali che passano per l'origine di cui conosciamo l'equazione.

Sia ora $[\phi_1, \phi_2] \in \mathbb{T}^2$ la proiezione del punto fisso (x_1, x_2) per l'azione di L in $\mathbb{R}^2$. Siano ℓ_u e ℓ_s delle rette in $\mathbb{R}^2$, parallele a W_s e W_u rispettivamente, che si intersecano in $[\phi_1, \phi_2]$. Denotiamo le proiezioni in $\mathbb{T}^2$ di queste rette nel modo seguente:

$$W^s[\phi_1, \phi_2] = \pi(\ell_s); \quad W^u[\phi_1, \phi_2] = \pi(\ell_u)$$

Facciamo ora vedere che queste proiezioni altro non sono che gli insiemi stabili e instabili associati a $[\phi_1, \phi_2]$ come punto fisso in $\mathbb{T}^2$ per la dinamica iperbolica. Si noti che enunciamo il risultato in generale anche se il solo punto fisso come sappiamo è l'origine. D'ora in poi

indicheremo λ_s l'autovalore corrispondente alla direzione stabile, $|\lambda_s| < 1$, e λ_u l'autovalore corrispondente alla direzione instabile, $|\lambda_u| > 1$.

Proposizione 5.21

1. $W^s[\phi_1, \phi_2]$ è l'insieme stabile di $[\phi_1, \phi_2]$, cioè se $[\phi'_1, \phi'_2] \in W^s[\phi_1, \phi_2]$ allora

$$\lim_{n \rightarrow \infty} d(L^n[\phi'_1, \phi'_2], [\phi_1, \phi_2]) = 0$$

dove d è la distanza in $\mathbb{T}^2$ definita dalla (4.32).

2. Analogamente, $W^u[\phi_1, \phi_2]$ è l'insieme instabile:

$$\lim_{n \rightarrow \infty} d(L^n[\phi'_1, \phi'_2], [\phi_1, \phi_2]) = +\infty$$

se $[\phi'_1, \phi'_2] \in W^u[\phi_1, \phi_2]$.

Dimostrazione. Cominciamo col punto 1. L'azione di L come trasformazione in $\mathbb{R}^2$ è la solita: $L(\mathbf{x}) = L \cdot \mathbf{x} \forall \mathbf{x} \in \mathbb{R}^2$. Supponiamo che i punti (x_1, x_2) e (x'_1, x'_2) di $\mathbb{R}^2$ giacciono su una retta parallela a W_s . Denotiamo σ il segmento che congiunge (x_1, x_2) e (x'_1, x'_2) . Per la linearità di L anche $L^n(\sigma)$ è un segmento di retta che rimane parallelo a W_s . Per di più, la lunghezza di $L^n(\sigma)$ vale chiaramente $\lambda_s^n l(\sigma)$ dato che la propagazione avviene parallelamente alla direzione stabile, lungo la quale L agisce come moltiplicazione per λ_s . Pertanto $\|(L^n(x'_1, x'_2), (x_1, x_2))\| \rightarrow 0$ per $n \rightarrow \infty$ e ciò implica $\lim_{n \rightarrow \infty} d(L^n[\phi'_1, \phi'_2], [\phi_1, \phi_2]) = 0$. Nel caso 2 il ragionamento è del tutto analogo, e la proposizione è dimostrata.

La compattificazione di $\mathbb{R}^2$ in $\mathbb{T}^2$ tramite la quozientazione fa sì che gli insiemi stabili e instabili siano molto più complicati che delle semplici rette come in $\mathbb{R}^2$. Si ha infatti

Proposizione 5.22 *Sia L una mappa iperbolica. Allora ambo gli insiemi $W^s[\phi_1, \phi_2]$ e $W^u[\phi_1, \phi_2]$ sono densi in $\mathbb{T}^2$ per ogni $[\phi_1, \phi_2] \in \mathbb{T}^2$.*

Dimostrazione. Vogliamo ridurci al teorema di Jacobi sulla densità dell'orbita delle rotazioni del cerchio di angolo irrazionale facendo anzitutto vedere che W_s e W_u sono rette a pendenza (=coefficiente angolare) irrazionale nel piano. Infatti se così non fosse, considerando ad esempio W_s , questo insieme dovrebbe necessariamente contenere un punto

a coordinate intere (N_1, N_2) . Dato che L ha coefficienti interi, tutti gli iterati di (N_1, N_2) saranno punti a coefficienti interi il che non è possibile perché $\|(L^n(N_1, N_2))\| \rightarrow 0$ per $n \rightarrow \infty$.

Cosideriamo ora le intersezioni successive di W_s con la retta $x_2 = N_2$ in $\mathbb{R}^2$. Sia $x_1(N_2)$ la coordinata ascissa di questa intersezione. Si noti che $x_1(N_2)$ è il reciproco del coefficiente angolare di W_s che è un numero irrazionale. Dunque $x_1(2) = 2x_1(1)$ e in generale $x_1(N) = Nx_1(1)$.

Ricordiamo ora che il punto di coordinate $(x_1(j), j)$ in $\mathbb{R}^2$, $j \in \mathbb{Z}$, si proietta sul punto $[\alpha_j, 0]$ su $\mathbb{T}^2$, $0 \leq \alpha_j \leq 1$. La retta $x_2 = 0$ definisce una circonferenza in $\mathbb{T}^2$, e i numeri α_j sono le immagini successive di 0 rispetto ad una rotazione irrazionale della circonferenza. Dunque per il teorema di Jacobi esse sono densi nella circonferenza. Ragionando allo stesso modo si trova che le proiezioni dei punti $(j, x_1(j))$ sono dense sulla circonferenza. Poiché $\mathbb{T}^2 = S^1 \times S^1$ la proposizione è dimostrata.

L'insieme stabile e l'insieme instabile attorno ad ogni punto del toro sono dunque l'immagine su $\mathbb{T}^2$ delle rette W_s e W_u in $\mathbb{R}^2$ tramite la proiezione π . Per la proposizione appena provata ambedue queste curve devono riempire densamente il toro medesimo. Poiché i coefficienti angolari di W_s e W_u sono differenti, anche le loro proiezioni dovranno necessariamente incontrarsi in un insieme denso sul toro. Questi punti di intersezione hanno una notevole importanza.

Definizione 5.39 Sia $[\phi_1, \phi_2] \in \mathbb{T}^2$ un punto periodico per L . Un punto qualsiasi $[p_1, p_2] \neq [\phi_1, \phi_2]$ che appartiene a $W_u[\phi_1, \phi_2] \cap W_s[\phi_1, \phi_2]$ si dice punto omoclinico.

Osserviamo che $W_u[\phi_1, \phi_2]$ e $W_s[\phi_1, \phi_2]$ si incontrano sempre formando un angolo non nullo ad un punto omoclinico. In tal caso il punto omoclinico si dice *trasverso*. Si può quindi concludere:

Proposizione 5.23 I punti omoclinici trasversi sono densi in $\mathbb{T}^2$.

Queste idee ci permettono di provare che la dinamica discreta è in questo caso caotica per ogni dato iniziale. Più precisamente si ha

Teorema 5.14 Sia L un automorfismo iperbolico del toro $\mathbb{T}^2$. Allora

1. I punti periodici di L sono densi in $\mathbb{T}^2$;

2. L è topologicamente transitivo;
3. L ha dipendenza delicata dalle condizioni iniziali.

Osservazione 5.58

La dinamica discreta generata dall'iterazione di L è pertanto caotica, secondo la definizione data nel Capitolo precedente.

Dimostrazione. Il punto 1 è la Proposizione 4.20 già dimostrata. Il punto 2 lo proveremo in Appendice. Proviamo qui il punto 3.

Infatti, se $\phi \in \mathbb{T}^2$ e $\chi \in W_u[\phi]$, allora ogni iterazione di L allunga la distanza fra ϕ e χ , almeno lungo W_u . Di conseguenza si ha dipendenza delicata per ogni condizione iniziale sul toro $\mathbb{T}^2$.

6 Partizione di Markov e dinamica simbolica

Abbiamo visto nella mappa quadratica che la descrizione miglior dell'evoluzione caotica si ottiene tramite la dinamica simbolica. La stessa cosa vale anche in questo esempio. Per generare la dinamica simbolica occorrerà qui introdurre una particolare partizione del toro in insiemi disgiunti nota come *partizione di Markov*². Per semplificare le cose trattiamo solo l'esempio particolare dell'automorfismo iperbolico generato dalla matrice $L \in \text{SL}(2, \mathbb{Z})$ definita da

$$L = \begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix}$$

Gli autovalori di L sono $\frac{1}{2}(1 \pm \sqrt{5})$. L'autovettore corrispondente all'autovalore instabile è la retta $y = \frac{1}{2}(-1 + \sqrt{5})x$ mentre l'autovettore corrispondente all'autovalore stabile è la retta $y = -\frac{1}{2}(1 + \sqrt{5})x$.

Esercizio 6.42

Provare queste affermazioni.

²A.A.Markov (Ryazan 1856-S.Pietroburgo 1922), matematico russo professore all'Università di S.Pietroburgo. Fu uno dei principali fondatori dell'odierno calcolo delle probabilità.

Per prima cosa costruiamo dei rettangoli i cui lati giacciono negli insiemi stabili e instabili di $[0]$. Più precisamente, consideriamo l'intervallo da a a b in $W_s[0]$ e l'intervallo da c a d in $W_u[0]$, come nella fig.4.1.

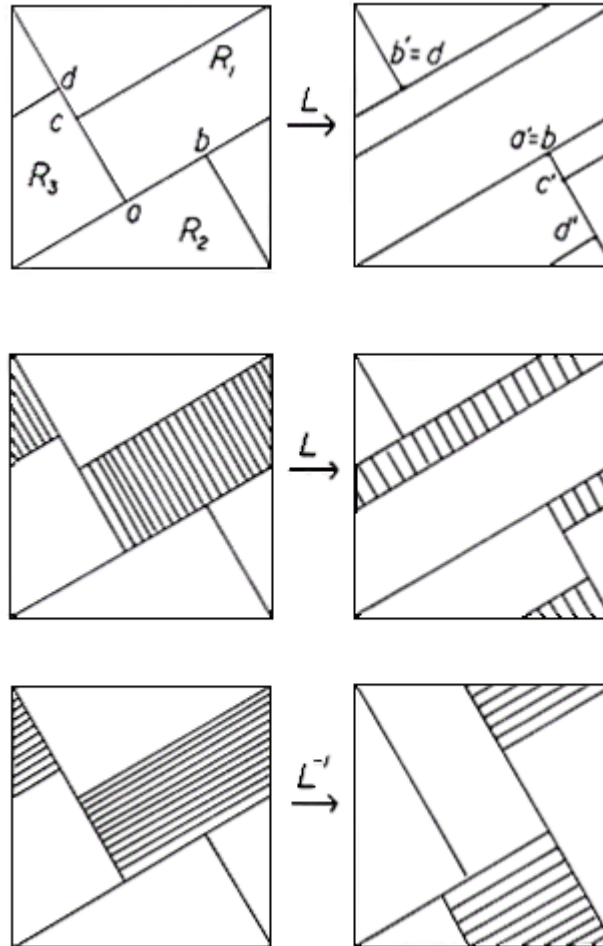


Figure 6.2: L'azione di L su una partizione di Markov

Allo stesso modo, i contorni instabili degli R_i stanno nell'intervallo fra c e d . Questo intervallo viene contratto da L^{-1} , cosicché l'intera orbita all'indietro dei punti nel contorno instabile giace in questo insieme (si veda la fig.4.1).

Per studiare la dinamica su $\mathbb{T}^2$, cominciamo col notare che ogniqualvolta $L(R_i)$ incontra l'interno di R_j l'immagine passa completamente attraverso R_j nella direzione instabile. Allo stesso modo ogniqualvolta $L^{-1}(R_i)$ incontra l'interno di R_j l'immagine passa completamente

attraverso R_j nella direzione stabile. I rettangoli che hanno questa proprietà e i cui lati stanno negli insiemi stabile e instabili formano per definizione una *partizione di Markov* per la mappa L . Questa partizione ci permetterà di costruire una dinamica simbolica, perché immagini in avanti degli R_i daranno sempre rettangoli "instabili" e immagini all'indietro rettangoli "stabili".

Si noti che $L(R_1)$ taglia l'interno sia di R_2 che di R_3 , $L(R_2)$ taglia l'interno sia di R_1 che di R_3 , e $L(R_3)$ taglia solo l'interno di R_2 . Per il momento ignoriamo il fatto che $L(R_1)$ incontra R_1 lungo la frontiera di R_1 e R_2 . Allo stesso modo, le intersezioni $L(R_2) \cap R_2$, $L(R_3) \cap R_1$ e $L(R_3) \cap R_2$ non sono vuote ma sono contenute nei contorni dei rettangoli. Così si può sperare di stabilire un'equivalenza con la sottotraslazione di tipo finito Σ_B dove la matrice di transizione è definita da

$$B = \begin{pmatrix} 0 & 1 & 1 \\ 1 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix}$$

Ci sono tuttavia molti problemi nello stabilire questa equivalenza. Per prima cosa, nessuno dei punti fissi $(\dots 111 \dots)$, $(\dots 222 \dots)$, $(\dots 333 \dots)$ sono successioni che non possono appartenere a Σ_B ; tuttavia sappiamo che in $\mathbb{T}^2$ c'è un punto fisso, e precisamente $[0]$. Per di più, c'è un'ambiguità nell'assegnazione delle successioni quando il punto o una delle sue immagini giace in uno dei contorni di un rettangolo.

Per porre rimedio a questi problemi lavoreremo con una sottotraslazione *quoziente*. Supponiamo che un punto p giaccia sul contorno stabile di $R_2 \cap R_3$. Sia $S(p) = (\dots s_0 s_1 s_2 \dots)$ una successione associata in modo naturale a p . Poiché $p \in R_2 \cap R_3$, dobbiamo avere o $s_0(p) = 2$ oppure $s_0(p) = 3$. Ora $L(p)$ giace nell'intersezione $R_2 \cap R_3$ come vediamo dalla fig.4.5.

Le immagini successive di p rimbalzano avanti e indietro fra $R_2 \cap R_3$ e $R_1 \cap R_3$. Torniamo ora a $S(p)$. Se $s_0(p) = 2$ allora $s_1(p)$ vale 1 o 2. Ma 2 non può seguire 2, e quindi $s_1(p) = 1$. Continuando questo ragionamento, dobbiamo avere $S(p) = (\dots s_{-1} \cdot 2121 \dots)$. D'altra parte, se $s_0(p) = 3$, allora dovrà essere $S(p) = (\dots, s_{-1} 3232 \dots)$. Questo significa che le due scelte possibili $(\dots, s_{-2}, s_{-1} 2121 \dots)$ e $(\dots, s_{-2}, s_{-1} 3232 \dots)$ devono rappresentare lo stesso punto in $\mathbb{T}^2$; in altre parole, queste successioni devono essere identificate. Più in generale, anche

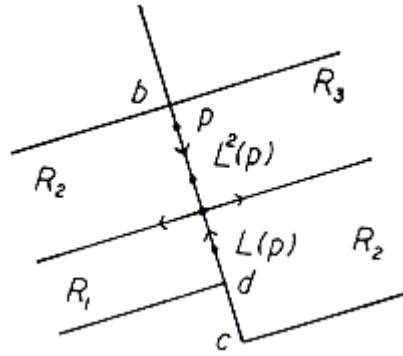


Figure 6.3:

tutte le successioni di una qualsiasi delle due forme

$$(\dots, s_{k-1}, s_k 2121 \dots)$$

$$(\dots, s_{k-1}, s_k 3232 \dots)$$

devono essere identificate, perché rappresentano punti che prima o poi vanno a finire su $W^s[0]$. lasciamo al lettore la verifica, tramite la figura 4.6, che la coppia di successioni della forma

$$(\dots, 1212, s_k, s_{k+1} \dots)$$

$$(\dots, 2121, s_k, s_{k+1} \dots)$$

devono essere identificate, così come le successioni della forma

$$(\dots, 2323, s_k, s_{k+1} \dots)$$

$$(\dots, 3232, s_k, s_{k+1} \dots)$$

Ora si denoti $\tilde{\Sigma}_B$ il "quoziente" della sottotraslazione di tipo finito ottenuto facendo tutte le identificazioni precedenti. Si noti che σ è definita in modo naturale su $\tilde{\Sigma}_B$. Lasciamo al lettore la verifica che che S definisce un'applicazione biunivoca di da $\mathbb{T}^2$ a $\tilde{\Sigma}_B$ che coniuga L con σ . In altre parole

$$S \circ L \circ S^{-1} = \sigma$$

Osservazione 6.59

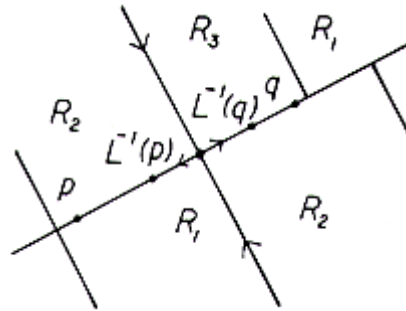


Figure 6.4:

1. Tutti i punti della forma $(\dots 1212\dots)$, $(\dots 2121\dots)$, $(\dots 2323\dots)$ e $(\dots 3232\dots)$ devono essere considerati identici in $\tilde{\Sigma}_B$. Poiché $\sigma(\dots 1212\dots) = (\dots 2121\dots)$ da ciò segue che queste sono le successioni che rappresentano il punto fisso a 0 .
2. Le sole identificazioni previste dal procedimento precedente avvengono su $W^s[0]$ e $W^u[0]$, dove la dinamica di L è relativamente semplice. Tutti gli altri punti periodici di L si trovano nel complemento di questi due insiemi. Come sappiamo, $W^s[0]$ e $W^u[0]$ sono densi in $\mathbb{T}^2$. Ciò nondimeno S si comporta bene sul complemento e descrive completamente la dinamica su questo insieme.
3. Non discuteremo la continuità di S poiché le identificazioni in $\tilde{\Sigma}_B$ significano che la topologia sullo spazio delle successioni è differente da quella solita.

La partizione di Markov costruita sopra è allo stesso tempo del tutto generale e del tutto particolare. È particolare perché gli elementi della partizione sono veri rettangoli; in realtà tutto ciò di cui abbiamo bisogno è che i contorni degli elementi della partizione appartengano agli appropriati inisemi stabili e instabili. D'altra parte l'uso degli insiemi stabili e instabili conosciuti, come sopra, nel costruire partizioni di Markov è un'operazione completamente generale. Tutto ciò che serve è che la mappa lasci invarianti questi insiemi.

Capitolo 7

Appendice

1 Dimostrazione del Teorema 6.7

Cominciamo col dimostrare che S è iniettiva. Siano $(x, y) \in \Lambda$ e supponiamo $S(x) = S(y)$. Dobbiamo far vedere che $x = y$. Se $S(x) = S(y)$ $f_\mu^n(x)$ e $f_\mu^n(y)$ stanno dalla stessa parte rispetto al punto medio $1/2$ di I . Ciò implica che $f_\mu(x)$ è monotona nell'intervallo di estremi $f_\mu^n(x)$ e $f_\mu^n(y)$. Di conseguenza tutti i punti di questo intervallo devono rimanere in $I_0 \cup I_1$ sotto l'azione della dinamica, ma ciò contraddice il fatto che Λ non contiene alcun intervallo.

Adesso facciamo vedere che S è su, cioè la sua immagine è Σ_2 . A questo scopo, introduciamo la notazione seguente: sia $J \subset I$ un intervallo chiuso. Sia

$$f_\mu^{-n}(J) := \{x \in I \mid f_\mu^n(x) \in J\}$$

In particolare, $f_\mu^{-1}(J)$ denota la preimmaagine di J rispetto a f_μ . Si osservi che se $J \subset I$ è un intervallo chiuso, allora $f_\mu^{-1}(J)$ consiste di due sottointervalli, uno in I_0 e l'altro in I_1 .

Sia ora $\mathbf{s} = (s_0, s_1, s_2, \dots)$. Dobbiamo costruire un $x \in \Lambda$ tale che $S(x) = \mathbf{s}$. A questo scopo definiamo

$$\begin{aligned} I_{s_0 s_1 \dots s_n} &= \{x \in I \mid x \in I_0, f_\mu(x) \in I_{s_1}, \dots, f_\mu^n(x) \in I_{s_n}\} \\ &= I_{s_0} \cap f_\mu^{-1}(I_{s_1}) \cap \dots \cap f_\mu^{-n}(I_{s_n}) \end{aligned}$$

Affermiamo che $I_{s_0 s_1 \dots s_n}$ forma una successione di intervalli chiusi e non vuoti in cui il successivo è sempre contenuto nel precedente. Si noti che

$$I_{s_0 s_1 \dots s_n} = I_{s_0} \cap f_\mu^{-1}(I_{s_1, \dots, s_n})$$

e pertanto possiamo concludere che l'insieme $\bigcap_{n \geq 0} I_{s_0 s_1 \dots s_n}$ non è vuoto. Si noti che se $x \in \bigcap_{n \geq 0} I_{s_0 s_1 \dots s_n}$ allora $x \in I_{s_0}$, $f_\mu(x) \in I_{s_0}$, ecc. Pertanto $S(x) = (s_0, s_1, s_2, \dots)$. Questo dimostra che S è suriettiva.

Si noti poi che $\bigcap_{n \geq 0} I_{s_0 s_1 \dots s_n}$ consiste di un unico punto. Questo segue subito dal fatto che S è $1-1$. In particolare abbiamo che la lunghezza di $I_{s_0 s_1 \dots s_n}$ tende a 0 per $n \rightarrow \infty$.

Dimostriamo ora la continuità di S . Sia $x \in \Lambda$, e supponiamo che sia $S(x) = (s_0, s_1, s_2, \dots)$. Sia $\epsilon > 0$, e scegliamo n in modo tale da avere $2^{-n} < \epsilon$. Consideriamo gli intervalli chiusi $I_{t_0 t_1 \dots t_n}$ definiti sopra per tutte le possibili combinazioni $t_0 t_1 \dots t_n$. Questi intervalli sono tutti disgiunti, e Λ è contenuto nella loro unione. Ci sono 2^{n+1} simili intervalli, e $I_{s_0 s_1 \dots s_n}$ è uno di essi. Pertanto possiamo scegliere δ tale che $|x - y| < \delta$, e $y \in \Lambda$ implica che $y \in I_{s_0 s_1 \dots s_n}$. Pertanto i primi $n + 1$ termini di $S(x)$ e $S(y)$ sono gli stessi. Pertanto, come sappiamo,

$$d[S(x), S(y)] < \frac{1}{2^n} < \epsilon$$

e questo prova la continuità di S . È agevole poi verificare che anche S^{-1} è continua. Pertanto S è un omeomorfismo.

Adesso dimostriamo che la codificazione generata da S rende equivalenti le dinamiche di f_μ su Λ e di σ su Σ_2 , cioè dimostriamo che $S \circ f_\mu = \sigma \circ S$. Un punto $x \in \Lambda$ è individuato univocamente dalla successione di intervalli inscatolati

$$\bigcap_{n \geq 0} I_{s_0 s_1 \dots s_n}$$

determinati dall'itinerario $S(x)$. Ora:

$$I_{s_0 s_1 \dots s_n} = I_{s_0} \cap f_\mu^{-1}(I_{s_1}) \cap \dots \cap f_\mu^{-n}(I_{s_n})$$

cosicché $f_\mu(I_{s_0 s_1 \dots s_n})$ può essere scritta nel modo seguente

$$I_{s_1} \cap f_\mu^{-1}(I_{s_2}) \cap \dots \cap f_\mu^{-n}(I_{s_n}) = I_{s_1 s_2 \dots s_n}$$

dato che $f_\mu(I_{s_0}) = I$. Pertanto

$$S f_\mu(x) = S f_\mu \left(\bigcap_{n \geq 0} I_{s_0 s_1 \dots s_n} \right) = S \left(\bigcap_{n=1}^{\infty} I_{s_0 s_1 \dots s_n} \right) = (s_1, s_2, \dots) = \sigma S(x)$$

e questo prova il teorema.

2 Dimostrazione dei Teoremi 2.9,2.10,2.11

Dimostrazione del Teorema 2.9. Si consideri la funzione di due variabili $G(x, \lambda) := f_\lambda(x) - x$, definita e regolare per $(x, \lambda) \in I \times N$. Per ipotesi $G(x_0, \lambda_0) = 0$. Inoltre:

$$\frac{\partial G}{\partial x}(x_0, \lambda_0) = f'_\lambda(x_0) - 1 \neq 0$$

Dunque possiamo applicare il teorema della funzione implicita e desumere l'esistenza di un intervallo I_1 attorno a x_0 , di un intervallo N_1 attorno a λ_0 e di una funzione regolare $p : N_1 \rightarrow I_1$ tale che $p(x_0) = \lambda_0$ e $G(p(\lambda), \lambda) = 0$ per ogni $\lambda \in N$. Per di più $G(x, \lambda) \neq 0$ se $x \neq p(\lambda)$, e ciò conclude la dimostrazione.

Dimostrazione del Teorema 2.10. Sia ancora $G(x, \lambda) := f_\lambda(x) - x$, e si noti che la condizione $G = 0$ significa che f_λ ha un punto fisso a x . Applichiamo ancora il teorema della funzione implicita a G . Possiamo farlo perché $G(0, \lambda_0) = 0$ e

$$\frac{\partial G}{\partial \lambda}(0, \lambda_0) = \frac{\partial f_\lambda}{\partial \lambda} \Big|_{\lambda=\lambda_0}(0) \neq 0$$

Dunque esisterà una funzione regolare $p(x)$ in un intorno di $x = 0$ tale che $G(x, p(x)) = 0$. Per la regola di derivazione delle funzioni composte possiamo scrivere

$$\frac{\partial G}{\partial x} + \frac{\partial G}{\partial \lambda} p'(x) = 0$$

Pertanto

$$p'(x) = -\frac{\partial G}{\partial x}(x, p(x)) / \frac{\partial G}{\partial \lambda}(x, p(x))$$

Derivando questa espressione si trova subito

$$p''(0) = -\frac{\frac{\partial^2 G}{\partial x^2}(0) \frac{\partial G}{\partial \lambda} \Big|_{\lambda=\lambda_0}(0)}{\left(\frac{\partial G}{\partial \lambda}\right)^2} = -\frac{f''_{\lambda_0}(0)}{\frac{\partial f}{\partial \lambda} \Big|_{\lambda=\lambda_0}(0)}$$

e ciò conclude la dimostrazione.

Dimostrazione del Teorema 2.11. Qui definiamo invece $G(x, \lambda) := f_\lambda^2(x) - x$. Non si può applicare direttamente il teorema della funzione implicita dato che

$$\begin{aligned} \frac{\partial G}{\partial \lambda}(0, \lambda_0) &= \frac{\partial f_\lambda}{\partial x}(0, \lambda_0) \frac{\partial f_\lambda}{\partial \lambda}(0, \lambda_0) + \frac{\partial f_\lambda}{\partial \lambda}(0, \lambda_0) = \\ &= \frac{\partial f_\lambda}{\partial \lambda}(0, \lambda_0) \left[1 + \frac{\partial f_\lambda}{\partial x}(0, \lambda_0) \right] = 0 \end{aligned}$$

per la Condizione 2. Definiamo allora

$$H(x; \lambda) = \begin{cases} \frac{G(x, \lambda)}{x}, & x \neq 0 \\ \frac{\partial G}{\partial x}(0, \lambda), & x = 0 \end{cases}$$

Si verifica subito che H è regolare e che

$$\begin{aligned} \frac{\partial H}{\partial x}(0, \lambda_0) &= \frac{1}{2} \frac{\partial^2 G}{\partial x^2}(0, \lambda_0) \\ \frac{\partial^2 H}{\partial x^2}(0, \lambda_0) &= \frac{1}{3} \frac{\partial^3 G}{\partial x^3}(0, \lambda_0) \end{aligned}$$

Possiamo allora applicare il teorema della funzione implicita a H . Si noti che

$$\begin{aligned} H(0, \lambda_0) &= \frac{\partial G}{\partial x}(0, \lambda_0) = \\ &= (f_{\lambda_0}^2)'(0) - 1 = f'_{\lambda_0}(0) \cdot f'_{\lambda_0}(0) - 1 = 0 \end{aligned}$$

D'altra parte per ipotesi si ha

$$\begin{aligned} \frac{\partial H}{\partial \lambda}(0, \lambda_0) &= \frac{\partial}{\partial \lambda} \Big|_{\lambda=\lambda_0} (0, \lambda_0)((f_{\lambda_0}^2)'(0) - 1) = \\ &= \frac{\partial (f_{\lambda}^2)'}{\partial \lambda}(0) \neq 0 \end{aligned}$$

Pertanto esiste una funzione regolare $p(x)$ definita in un intorno di 0 che soddisfa $p(0) = \lambda_0$ e $H(x, p(x)) = 0$. In particolare

$$\frac{1}{x} G(x, p(x)) = 0$$

per $x \neq 0$. Ne segue che x è un punto periodico di periodo 2 per $f_{p(x)}$. Si noti che x non è fissato da $f_{p(x)}$ per il teorema 2.9. Come sopra, possiamo calcolare $p'(0)$:

$$p'(0) = - \frac{\frac{\partial H}{\partial x}(0, \lambda_0)}{\frac{\partial H}{\partial \lambda}(0, \lambda_0)} = 0$$

dato che

$$(f_{\lambda_0}^2)''(0) = f''_{\lambda_0}(0) \cdot (f'_{\lambda_0}(0))^2 + f''_{\lambda_0}(0) \cdot f'_{\lambda_0}(0) = 0$$

dove abbiamo usato l'ipotesi $f'_{\lambda_0}(0) = -1$. Ciò completa la dimostrazione.

3 Dimostrazione delle proposizioni sulla derivata di Schwarz

Dimostrazione della Proposizione 3.11.

1. Se $P'(x)$ ha N radici reali e distinte, denotate $a_i : i = 1, \dots, N$ possiamo scrivere

$$P'(x) = \prod_{i=1}^n (x - a_i), \quad a_1 < a_2 < \dots < a_N$$

Pertanto si ha

$$\begin{aligned} P''(x) &= \sum_{j=1}^N \frac{P'(x)}{x - a_j} = \sum_{j=1}^N \frac{\prod_{i=1}^n (x - a_i)}{x - a_j} \\ P'''(x) &= \sum_{j=1}^N \sum_{\substack{k=1 \\ k \neq j}}^N \frac{\prod_{i=1}^n (x - a_i)}{(x - a_j)(x - a_k)} \end{aligned}$$

Di conseguenza possiamo scrivere

$$\begin{aligned} SP(x) &= \sum_{j \neq k} \frac{1}{(x - a_j)(x - a_k)} - \frac{3}{2} \left(\sum_{j=1}^N \frac{1}{x - a_j} \right)^2 \\ &= -\frac{1}{2} \left(\sum_{j=1}^N \frac{1}{x - a_j} \right)^2 - \left(\sum_{j=1}^N \frac{1}{x - a_j} \right)^2 < 0 \end{aligned}$$

2. Facendo uso della regola di derivazione delle funzioni composte si ha

$$(f \circ g)''(x) = f''(g(x)) \cdot (g'(x))^2 + f'(g(x)) \cdot g''(x)$$

e inoltre

$$(f \circ g)'''(x) = f'''(g(x)) \cdot (g'(x))^3 + 3f''(g(x)) \cdot g''(x) \cdot g'(x) + f'(g(x)) \cdot g'''(x)$$

Ne segue che

$$S(f \circ g)(x) = Sf(g(x)) \cdot (g'(x))^2 + Sg(x)$$

e quindi $S(f \circ g)(x) < 0$.

Per dimostrare il punto 3 dobbiamo premettere tre lemmi.

Lemma 3.2 *Se $Sf < 0$ allora $f'(x)$ non può avere un minimo locale positivo o un massimo locale negativo.*

Dimostrazione. Sia x_0 un punto critico di $f'(x)$, cioè sia $f''(x_0) = 0$. Poiché $Sf < 0$, risulterà $f'''(x_0)/f'(x_0) < 0$ e quindi $f'''(x_0)$ e $f'(x_0)$ hanno segni opposti, e questo chiaramente basta a provare l'asserzione: ad esempio, se $f'(x)$ avesse un minimo locale positivo in x_0 la sua derivata seconda ivi, cioè $f'''(x_0)$, dovrebbe essere positiva.

Questo lemma implica che il grafico di $f'(x)$ deve tagliare l'asse delle x fra due punti critici successivi. In particolare, quindi, deve esistere un punto critico di f fra questi due punti.

Lemma 3.3 *Se $f(x)$ ha un numero finito di punti critici, $f^m(x)$ gode della stessa proprietà.*

Dimostrazione. Siac arbitrario nel codominio di f . Allora la preimmagine $f^{-1}(c)$ è formata da un numero finito di punti. Infatti fra due punti qualsiasi della premimmagine di c deve esserci almeno un punto critico di f . Ne segue senza difficoltà che anche la controimmagine $f^{-m}(c) := \{x \mid f^m(x) = c\}$ è un insieme finito.

Ora supponiamo $(f^m)'(x) = 0$. per la regola di derivazione delle funzioni composte possiamo scrivere

$$(f^m)'(x) = \prod_{i=1}^{m-1} f'(f^i(x))$$

Pertanto $f^i(x)$ è un punto critico di f per almeno un i , $0 \leq i \leq m-1$. Quindi l'insieme dei punti critici di f^m è composto dall'unione delle controimmagini di ordine minore di m dell'insieme dei punti critici di m assieme alle loro orbite. Per l'osservazione precedente, anche questo insieme di punti è finito.

Lemma 3.4 *Supponiamo che $f(x)$ abbia un numero finito di punti critici e che $Sf < 0$. Allora per ogni intero m f ha solo un numero finito di punti periodici di periodo m .*

Dimostrazione. Sia $g := f^m$. Allora i punti periodici di periodo m di f sono punti fissi di g . Per la Proposizione 3.11, $Sg < 0$.

Supponiamo che g abbia infiniti punti fissi. Per il teorema del valor medio, esistono allora infiniti punti per cui $g'(x) = 1$. Fra ogni tre punti consecutivi per cui $g'(x) = 1$ ci deve essere un punto per cui $g'(x) < 1$. Infatti $g'(x) = 1$ non vale 1 identicamente su un qualche intervallo perché altrimenti risulterebbe $Sg = 0$ contraddicendo l'ipotesi $Sg < 0$. Inoltre, per il Lemma 3.2, $g'(x)$ non può avere un minimo locale positivo fra questi tre punti. Pertanto devono esserci punti per cui $g'(x) < 0$. Di conseguenza ci sono punti in cui $g'(x) = 0$. Questo

però implica che g ha infiniti punti critici. Questo contraddice il Lemma 3.3 e completa la prova del Lemma presente.

Conclusione della dimostrazione del Teorema 3.12. Sia p un punto periodico attrattivo di f di periodo m . Sia $W(p)$ l'intervallo massimale attorno a p tutti i punti del quale tendono a p asintoticamente tramite la mappa f^m (il cosiddetto *bacino di attrazione* di p). In formule, $W(p)$ è la componente connessa che contiene p dell'insieme $\{x \mid f^mj \rightarrow p \mid j \rightarrow \infty\}$. Chiaramente $W(p)$ è un intervallo aperto invariante rispetto a f^m , cioè $f^{om}W(p) \subset W(p)$. Supponiamo per il momento che p sia un punto fisso. Poiché $W(p) \subset W(p)$ e $W(p)$ è massimale, ne segue che o f conserva gli estremi (ℓ, r) di $W(p)$ oppure almeno uno fra ℓ e r è infinito. Nel caso finito si presentano tre possibilità:

1. $f(\ell) = \ell$ e $f(r) = r$;
2. $f(\ell) = r$ e $f(r) = \ell$;
3. $f(\ell) = f(r)$.

Se $f(\ell) = \ell$ e $f(r) = r$ allora il grafico di f mostra che esistono a, b che soddisfano le disuguaglianze $\ell < a < p < b < r$ e $f'(a) = f'(b) = 1$. Poiché $f'(p) < 1$ e $f'(x)$ non può avere un minimo locale positivo per il Lemma 3.2, ne segue che deve esistere un punto critico nell'intervallo $]a, b[$. Il secondo caso segue allo stesso modo considerando f^2 invece di f . Nel caso 3, f deve avere un minimo o un massimo fra ℓ e r , cosicché esiste un punto critico di f anche in questo caso. Il ragionamento però non è applicabile al caso in cui uno almeno fra ℓ e r è infinito. Tuttavia questi casi aggiungono al più due punti fissi.

Se p è un punto periodico, i medesimi ragionamenti mostrano l'esistenza di un punto critico di f^m in $W(p)$. Ora un punto nell'orbita di questo punto critico deve però essere un punto critico di f per la regola di derivazione delle funzioni composte. Ciò conclude la dimostrazione del teorema.

4 Completamento della dimostrazione del Teorema 5.14

L'unica cosa che rimande da provare è l'affermazione 2, cioè la transitività topologica. Siano U e V due aperti in $\mathbb{T}^2$. Ciò significa, per definizione, che le loro preimmagini attraverso la

proiezione π sono degli aperti di $\mathbb{R}^2$. Costruiamo ora esplicitamente un punto $\phi \in U$ ed un intero k tali che $L^k[\phi] \in V$. Possiamo, data la densità, scegliere punti $[\rho] \in U$ e $[\sigma] \in V$ omoclinici a $[0]$.

Sia ora $\epsilon > 0$. Consideriamo un intervallo aperto I_u di lunghezza $\delta > 0$ in $W_u[0]$ che contiene $[\sigma]$. L^n espande I_u di un fattore $|\lambda_u|^n$ e L^{-n} espande I_s del medesimo fattore. Ora si prenda n sufficientemente grande in modo che

1. $d(L^n[\rho], [0]) < \frac{\epsilon}{2}$;
2. $d(L^{-n}[\sigma], [0]) < \frac{\epsilon}{2}$;
3. $|\lambda_u|^n \delta > \epsilon$.

Qui al solito d è la distanza euclidea (in un intorno di $[0]$ indotta dalla distanza euclidea su $\mathbb{R}^2$ dalla proiezione π). Poiché $L^n(I_u)$ e $L^{-n}(I_s)$ sono paralleli a $W_u[0]$ e $W_s[0]$, rispettivamente, ne segue che $L^n(I_u) \cap L^{-n}(I_s) \neq \emptyset$. Sia $[\mathbf{q}]$ un punto in tale intersezione. Allora $[\mathbf{p}] := L^{-n}[\mathbf{q}] \in U$ e $L^n[\mathbf{q}] \in V$. Di conseguenza $L^{2n}[\mathbf{p}] \in V$, e si ottiene così il punto voluto.