

Corso di Laurea in Fisica

Equazioni Differenziali e Sistemi Dinamici

Giulio Schimperna

Dipartimento di Matematica, Università di Pavia

Via Ferrata 1, 27100 – PAVIA

E-mail: `giulio@dimat.unipv.it`

Homepage: `http://www-dimat.unipv.it/~giulio/sistdin05.html`

Versione del **18 Gennaio 2006**

Contenuto

1	Teoria generale	2
1.1	Richiami sugli spazi vettoriali normati	3
1.2	Esistenza e unicità in piccolo	6
1.3	Prolungamento ed esistenza in grande	10
1.4	Dipendenza continua dai dati	14
1.5	Alcuni tipi di equazioni	17
1.6	Equazioni autonome – rappresentazione delle soluzioni	20
2	Sistemi lineari	23
2.1	Richiami di algebra lineare – diagonalizzazione	23
2.2	Richiami di algebra lineare – forme canoniche	28
2.3	Teoria generale	33
2.4	Sistemi a coefficienti costanti diagonalizzabili	39
2.5	Sistemi a coefficienti costanti non diagonalizzabili	42
2.6	Equazioni lineari a coefficienti costanti di ordine superiore al primo	46
3	Stabilità	48
3.1	Stabilità dei sistemi lineari	49
3.2	Stabilità dei sistemi non lineari	58
3.3	Stabilità secondo Liapounov	65
4	Sistemi dinamici, orbite, attrattori	71
4.1	Un approccio astratto	71
4.2	Dissipatività e attrattori	78
5	Il teorema di Peano	88
5.1	Spazi metrici, completezza, compattezza	88
5.2	Il teorema di Ascoli	90
5.3	Esistenza senza unicità	94
5.4	Il “pennello di Peano”	97

1 Teoria generale

Questo primo capitolo comprende la parte più generale ed importante della teoria delle equazioni differenziali ordinarie (EDO nel seguito). Le notazioni introdotte e la teoria qui sviluppata saranno diffusamente utilizzate per tutto il resto del corso. Poiché una buona parte del materiale di questo capitolo è trattata sul testo [4], che dovrebbe essere già in possesso degli studenti, ci si limiterà a considerare gli argomenti che non sono sviluppati in [4]. Negli altri casi si darà solo qualche richiamo o la relativa indicazione bibliografica.

1.1 Richiami sugli spazi vettoriali normati

Per quanto riguarda la definizione di norma, di spazio vettoriale normato (s.v.n. nel seguito), di spazio di Banach e di spazio di Hilbert, si veda [4, §12.1, 12.2] (tralasciando il Teorema 12.2.3). Per indicare uno s.v.n. si useranno indifferentemente le notazioni X , ovvero $(X, \|\cdot\|_X)$, se vorremo enfatizzare quale è la norma di X . Ricordiamo che la norma è una funzione continua da X a $[0, +\infty)$; grazie alla disuguaglianza triangolare si ha infatti:

$$\| \|x\| - \|y\| \| \leq \|x - y\| \quad \forall x, y \in X. \quad (1.1)$$

Osserviamo ora che, dato un sottoinsieme C di uno s.v.n. V , la definizione [4, 12.2.1] di *successione di Cauchy* e la condizione [4, 12.2.2] di *completezza* si adattano senza difficoltà al caso in cui V è sostituito da C (che potrebbe *non* essere un sottospazio). In particolare, C si dice *completo* se e solo se ogni successione di Cauchy $\{x_n\}$ a valori in C ammette un limite x *appartenente* a C .

Esercizio 1.1. Mostrare che un sottoinsieme C di uno spazio di Banach V è completo se e solo se è chiuso. Analogamente, mostrare che ogni sottoinsieme completo C di un qualsiasi s.v.n. V è chiuso.

Sia ora data una successione $\{x_n\}$ a valori in uno spazio di Banach $(X, \|\cdot\|_X)$ e consideriamo la serie $\sum x_n$. Una utile condizione sufficiente perché questa converga nella norma di X è data dal

Teorema 1.2. *Vale la seguente implicazione:*

$$\sum \|x_n\|_X \text{ converge in } \mathbb{R} \implies \exists x \in X \text{ tale che } \sum_{n=0}^{\infty} x_n = x; \quad (1.2)$$

sotto tale condizione si ha inoltre che

$$\|x\|_X \leq \sum_{n=0}^{\infty} \|x_n\|_X. \quad (1.3)$$

La prova del Teorema non è difficile: sfruttando la disuguaglianza triangolare si mostra che la successione delle ridotte è di Cauchy per la norma di X ; quindi si sfrutta la completezza. Invitiamo il lettore a esplicitare i dettagli. Notiamo che questo risultato generalizza in senso astratto [4, Teorema X.2.11]. L'ipotesi in (1.2) viene detta *convergenza totale* della serie $\sum x_n$.

Vediamo ora di introdurre alcuni particolari spazi di Banach. Dato un sottoinsieme chiuso e limitato $K \subset \mathbb{R}^N$, consideriamo innanzitutto lo spazio vettoriale

$$C^0(K; \mathbb{R}^M) := \{f : K \rightarrow \mathbb{R}^M : f \text{ continua}\}. \quad (1.4)$$

Poniamo su tale spazio la norma

$$\|f\|_{\infty} := \sup_{x \in K} |f(x)|, \quad (1.5)$$

ove evidentemente il sup è finito; anzi, è più precisamente un massimo grazie al teorema di Weierstrass. Utilizzando le note proprietà del modulo in $\mathbb{R}^M$ è facile dimostrare che

$\|\cdot\|_\infty$ è effettivamente una norma. Si verifica inoltre che la convergenza nella norma $\|\cdot\|_\infty$ equivale al concetto di *convergenza uniforme* per successioni di funzioni [4, §X.2]. Vale inoltre la seguente

Proposizione 1.3. $C^0(K; \mathbb{R}^M)$ è completo, è cioè uno spazio di Banach.

Prova. Sia $\{f_n\}$ di Cauchy rispetto a $\|\cdot\|_\infty$. Di conseguenza, $\forall x \in K$, si ha che $\{f_n(x)\}$ è di Cauchy rispetto al modulo in $\mathbb{R}^M$. Fissiamo $\varepsilon > 0$. Poiché $\mathbb{R}^M$ è completo, $\{f_n(x)\}$ ammette un limite in $\mathbb{R}^M$, che battezziamo $f(x)$. Esplicitando la definizione della norma $\|\cdot\|_\infty$ nella condizione di Cauchy, quanto otteniamo è che $\exists \bar{n} \in \mathbb{N}$ dipendente da ε e tale che

$$|f_n(x) - f_m(x)| < \varepsilon \quad \forall n, m \geq \bar{n}, \quad \forall x \in K.$$

Ora, tenendo fisso n e facendo tendere m all'infinito, si ha subito

$$|f_n(x) - f(x)| \leq \varepsilon \quad \forall n \geq \bar{n}, \quad \forall x \in K,$$

che non è altro che la condizione di convergenza nella norma $\|\cdot\|_\infty$. Resta da dimostrare che f è continua e per questo si rimanda alla dimostrazione del Teorema 10.2.4 di [4] (conservazione della continuità per successioni di funzioni continue uniformemente convergenti). ■

Rinviamo al corso di “Complementi di Analisi Matematica di Base” per maggiori dettagli sulle proprietà della convergenza uniforme. Ci limitiamo ad osservare che, nel caso K non sia compatto, potrebbe accadere che il sup nella definizione (1.5) sia $+\infty$. In effetti, nel caso $A \subset \mathbb{R}^N$ sia un insieme qualunque, non è più possibile definire una struttura di spazio di Banach su $C^0(A; \mathbb{R}^M)$ utilizzando la convergenza uniforme su tutto A . Un modo per superare questo inconveniente è quello di considerare uno spazio più piccolo, ove la limitatezza è imposta a priori:

$$BC^0(A; \mathbb{R}^M) := \{f : A \rightarrow \mathbb{R}^M : f \text{ continua e limitata}\}, \quad (1.6)$$

Poniamo su tale spazio la norma

$$\|f\|_\infty := \sup_{x \in A} |f(x)|. \quad (1.7)$$

Si dimostra esattamente come nel caso precedente che lo spazio $BC^0(A)$ è di Banach rispetto alla norma $\|\cdot\|_\infty$.

Osservazione 1.4. Osserviamo che la norma $\|\cdot\|_\infty$ non è generata da un prodotto scalare; pertanto $C^0(K; \mathbb{R}^M)$ e $BC^0(A; \mathbb{R}^M)$ non sono spazi di Hilbert.

Osservazione 1.5. È anche interessante notare che, se A è un aperto di $\mathbb{R}^N$ è possibile definire una soddisfacente nozione di convergenza nello spazio vettoriale $C^0(A; \mathbb{R}^M)$ delle funzioni continue, non necessariamente limitate, da A a valori in $\mathbb{R}^M$. Tuttavia tale convergenza, di cui parleremo in dettaglio nel paragrafo § 1.4, non è indotta da alcuna norma.

Il seguente risultato costituisce il mattone fondamentale su cui si basano i risultati di esistenza ed unicità per le soluzioni di EDO:

Teorema 1.6. (Teorema delle contrazioni.) *Sia $(X, \|\cdot\|)$ uno spazio di Banach (o un suo sottoinsieme chiuso). Sia $T : X \rightarrow X$ una contrazione, ossia un'applicazione tale che*

$$\exists L < 1 : \quad \|Tx - Ty\| \leq L\|x - y\| \quad \forall x, y \in X. \quad (1.8)$$

Allora esiste uno ed un solo punto fisso di T , ossia

$$\exists! \bar{x} \in X \quad \text{tale che} \quad T\bar{x} = \bar{x}. \quad (1.9)$$

Prova. Comincio col mostrare l'unicità. Siano $\bar{x}_1, \bar{x}_2$ due punti fissi di T . Si ha allora

$$\|\bar{x}_1 - \bar{x}_2\| = \|T\bar{x}_1 - T\bar{x}_2\| \leq L\|\bar{x}_1 - \bar{x}_2\|,$$

da cui $\|\bar{x}_1 - \bar{x}_2\| = 0$ poichè $L < 1$. Dalle proprietà della norma segue che $\bar{x}_1 = \bar{x}_2$.

Per mostrare l'esistenza di $\bar{x}$, scelgo un *arbitrario* punto $x_0 \in X$ e pongo, per induzione, $x_j := Tx_{j-1}$, $j \in \mathbb{N}$, definendo in questo modo una successione $\{x_j\}$ di elementi di X . Sfruttando ripetutamente la (1.8), è facile mostrare che

$$\|x_j - x_{j-1}\| \leq L\|x_{j-1} - x_{j-2}\| \leq \cdots \leq L^{j-1}\|x_1 - x_0\|, \quad \forall j \geq 1. \quad (1.10)$$

Siano dati allora $n, p \in \mathbb{N} \setminus \{0\}$. Grazie alla disuguaglianza triangolare, si ha

$$\begin{aligned} \|x_{n+p} - x_n\| &\leq \sum_{j=n+1}^{n+p} \|x_j - x_{j-1}\| \\ &\leq \sum_{j=n+1}^{n+p} L^{j-1} \|x_1 - x_0\| \\ &= \text{(ponendo } k := j - n) \quad \|x_1 - x_0\| L^n \sum_{k=1}^p L^{k-1} \\ &= \|x_1 - x_0\| L^n \frac{1 - L^p}{1 - L} \leq \|x_1 - x_0\| \frac{L^n}{1 - L}. \end{aligned} \quad (1.11)$$

Poichè l'ultimo termine è infinitesimo per $n \rightarrow \infty$, si ha che la successione $\{x_j\}$ è di Cauchy e dunque converge a un limite, che battezzo $\bar{x}$. Passando al limite nella relazione (1.10), si ha subito che

$$0 = \lim_{j \rightarrow \infty} \|x_j - x_{j-1}\| = \lim_{j \rightarrow \infty} \|Tx_{j-1} - x_{j-1}\| = \|T\bar{x} - \bar{x}\|, \quad (1.12)$$

grazie alla continuità della norma (cf. (1.1)) e della funzione T . ■

Dati due spazi normati $(X, \|\cdot\|_X)$ e $(Y, \|\cdot\|_Y)$, si può dare in modo naturale la definizione di funzione *Lipschitziana* da X in Y . In particolare, una contrazione non è altro che una funzione di X in Y che verifica una condizione di Lipschitz con costante *strettamente* più piccola di 1. Le funzioni Lipschitziane verificano un'interessante proprietà di prolungabilità:

Proposizione 1.7. *Siano X, Y spazi di Banach, $A \subset X$ un sottoinsieme qualunque e sia $T : A \rightarrow Y$ Lipschitziana di costante L . Allora esiste, ed è unico, un prolungamento $\bar{T} : \bar{A} \rightarrow Y$ Lipschitziano di costante L .*

Prova. La prova dell'unicità di $\bar{T}$ è semplice e viene lasciata come esercizio. Per quanto riguarda l'esistenza, supponiamo di avere $\bar{x} \in \bar{A}$ e definiamo in modo opportuno $\bar{T}(\bar{x})$. Se $\bar{x} \in A$, poniamo semplicemente $\bar{T}(\bar{x}) := T(\bar{x})$. Altrimenti, consideriamo un'arbitraria successione $\{x_n\} \subset A$ tale che $x_n \rightarrow \bar{x}$. Tale successione esiste sicuramente, dato che $\bar{x} \in \bar{A}$. Ho allora

$$\|Tx_n - Tx_m\|_Y \leq L\|x_n - x_m\|_X \quad \forall n, m \in \mathbb{N};$$

dunque $\{Tx_n\}$ è di Cauchy in Y , poiché $\{x_n\}$ lo è in X . Essendo Y completo, esiste il limite di Tx_n , che battezzo $\bar{T}(\bar{x})$. Restano da fare due verifiche, ossia:

- che il valore $\bar{T}(\bar{x})$ non dipende dalla scelta della successione $\{x_n\}$;
- che la funzione $\bar{T}$ così definita è Lipschitziana di costante L .

Entrambi questi controlli vengono lasciati come (facile) esercizio. ■

Esercizio 1.8. Dimostrare che un analogo della Proposizione precedente vale anche nel caso in cui si supponga T uniformemente continua anziché Lipschitziana. Dedurre che la funzione $f : (0, 1) \rightarrow \mathbb{R}$, $f(x) = \sin(1/x)$ non è uniformemente continua.

1.2 Esistenza e unicità in piccolo

In questo paragrafo presentiamo il risultato più importante nella teoria generale delle EDO, che permette di individuare un'ampia classe di equazioni sicuramente “risolvibili”. Cominciamo introducendo una serie di notazioni fondamentali, che cercheremo di mantenere per tutto il corso della dispensa.

Sia Ω un sottoinsieme aperto di $\mathbb{R}^{N+1}$, $\mathbf{f} : \Omega \rightarrow \mathbb{R}^N$ una funzione nota che supporremo **sempre almeno continua**. Indicheremo in generale con $(t, \mathbf{x})$ la variabile in Ω , ove t è uno scalare e $\mathbf{x}$ è un vettore di $\mathbb{R}^N$. Consideriamo l'equazione differenziale

$$\mathbf{y}'(t) = \mathbf{f}(t, \mathbf{y}(t)), \quad (1.13)$$

la cui incognita è la funzione $\mathbf{y} = \mathbf{y}(t)$ definita su un opportuno sottoinsieme di $\mathbb{R}$ ed a valori in $\mathbb{R}^N$. È ben noto che, già nel caso in cui il secondo membro non dipende esplicitamente da $\mathbf{y}$, ossia siamo in presenza dell'equazione

$$\mathbf{y}'(t) = \mathbf{g}(t), \quad \text{ove } \mathbf{g} : (a, b) \rightarrow \mathbb{R}^N, \quad (1.14)$$

la soluzione non può essere unica; infatti ogni primitiva di $\mathbf{g}$ verifica la (1.14) ed esistono infinite primitive, due qualunque delle quali differiscono per una costante (vettoriale). Per “rendere unica” la soluzione, il modo canonico è naturalmente quello di imporre il valore della $\mathbf{y}$ in un arbitrario punto $t_0 \in (a, b)$, richiedendo cioè, oltre alla (1.14), che sia

$$\mathbf{y}(t_0) = \mathbf{y}_0, \quad \text{ove } \mathbf{y}_0 \in \mathbb{R}^N \text{ è un assegnato valore.} \quad (1.15)$$

Osserviamo peraltro che la continuità della $\mathbf{g}$ è sufficiente a garantire l'esistenza di una soluzione della (1.14) definita su tutto (a, b) , anche se naturalmente può succedere che $\mathbf{y}(t)$ esploda per $t \rightarrow b^-$ o per $t \rightarrow a^+$. Si noti infine che la determinazione esplicita delle soluzioni di (1.14) può essere difficile o impossibile anche per dati $\mathbf{g}$ dall'espressione analitica piuttosto semplice.

Tornando all'equazione (1.13), è abbastanza ovvio aspettarsi che questa conservi tutte le difficoltà connesse con la risoluzione di (1.14). Il punto dolente è che ne comporta anche di nuove, e piuttosto serie. Per affrontarle e, possibilmente, venirne a capo, cominciamo col provare a togliere di mezzo il problema della non-unicità. Scimmiettando (1.15), viene in mente di accoppiare (1.13) ad una condizione analoga. L'oggetto che otteniamo si chiama "Problema di Cauchy" o "ai valori iniziali" per l'equazione (1.13), e si formula come segue:

$$\begin{cases} \mathbf{y}'(t) = \mathbf{f}(t, \mathbf{y}(t)) \\ \mathbf{y}(t_0) = \mathbf{y}_0. \end{cases} \quad (1.16)$$

Tuttavia c'è un vincolo nuovo sul dato $(t_0, \mathbf{y}_0)$: dobbiamo senz'altro imporre che sia $(t_0, \mathbf{y}_0) \in \Omega$, ove, ricordiamo, Ω è il dominio di $\mathbf{f}$; altrimenti, (1.16) non avrebbe senso.

Notiamo inoltre che Ω , che abbiamo supposto aperto, può avere una forma anche molto irregolare; dunque, in generale, non esiste alcun intervallo privilegiato (a, b) , da richiedere come dominio della soluzione $\mathbf{y}$. In effetti, diverse ipotesi sul dominio di $\mathbf{y}$ ci porteranno a definire diverse nozioni di soluzione per (1.16), come vedremo nel corso del capitolo. Notiamo peraltro che almeno due condizioni sono imprescindibili: primo, affinché abbia senso la derivata in (1.13) dobbiamo comunque chiedere che il dominio di $\mathbf{y}$ sia un insieme aperto; secondo, dobbiamo supporre che contenga t_0 , altrimenti non avrebbe senso la condizione iniziale. Infine, aggiungiamo una terza richiesta universale, e cioè che il dominio di $\mathbf{y}$ sia un intervallo. Questa condizione, che sarà ulteriormente chiarita dalla teoria successiva, si può per il momento giustificare col fatto che ogni eventuale soluzione di (1.16) deve in ogni punto del suo dominio "tener conto" della condizione iniziale. Questo non può succedere se il dominio di $\mathbf{y}$ è sconnesso.

Tali richieste, da sole, portano alla più debole delle definizioni di soluzione per (1.16), quella di *soluzione in piccolo*, che presentiamo nel seguente modo un po' formale (ma comodo):

Definizione 1.9. Una soluzione in piccolo di (1.16) è una coppia $(\mathbf{y}, (a, b))$, ove (a, b) è un intervallo aperto contenente t_0 e la funzione $\mathbf{y} \in C^1((a, b); \mathbb{R}^N)$ verifica l'equazione (1.13) in ogni istante $t \in (a, b)$ e soddisfa la condizione iniziale $\mathbf{y}(t_0) = \mathbf{y}_0$.

Osservazione 1.10. Notiamo che, se $(\mathbf{y}, (a, b))$ è una soluzione, deve allora necessariamente aversi $\text{graf } \mathbf{y} \subset \Omega$, altrimenti l'equazione perderebbe significato. Ne segue in particolare che, se Ω non è connesso, ogni soluzione in piccolo "vive" nella componente connessa in cui è scelto il dato $(t_0, \mathbf{y}_0)$ e non può uscire da questa.

Dunque, la soluzione viene vista come una coppia formata dalla funzione $\mathbf{y}$ e dal suo dominio (a, b) . Naturalmente, se un'impostazione così formale dà fastidio (e se non c'è possibilità di confusione), si può continuare a pensare a una soluzione di (1.16)

come alla sola funzione $\mathbf{y}$. Tuttavia, sarà utile mantenere questo formalismo almeno per un po'.

Notiamo anche che la scelta della lettera t per indicare la variabile indipendente non è casuale; specialmente nella seconda parte del corso, infatti, ci abitueremo a pensare a t come a una variabile tempo e a $\mathbf{y}$ come una quantità che descrive qualche processo dinamico (si pensi ad esempio al moto di una particella: $\mathbf{y}(t)$ potrebbe ad esempio indicare la posizione in $\mathbb{R}^3$). In quest'ottica è chiaro perché spesso la condizione $\mathbf{y}(t_0) = \mathbf{y}_0$ è detta *condizione iniziale*. Anzi, in molti casi supporremo che l'“origine dei tempi” sia fissata in $t_0 = 0$. Sempre sotto questo punto di vista, in certe situazioni può essere utile andare a considerare l'evoluzione della variabile $\mathbf{y}$ soltanto *a partire da* t_0 . Questo porta ad introdurre il cosiddetto problema di Cauchy *in avanti*

$$\begin{cases} \mathbf{y}'(t) = \mathbf{f}(t, \mathbf{y}(t)) & \text{per } t \geq t_0 \\ \mathbf{y}(t_0) = \mathbf{y}_0. \end{cases} \quad (1.17)$$

È abbastanza naturale adattare la definizione di soluzione in piccolo come segue:

Definizione 1.11. Una soluzione in piccolo di (1.17), è una coppia $(\mathbf{y}, T)$, ove $T > t_0$ è detto tempo finale della soluzione e $\mathbf{y} \in C^1([t_0, T]; \mathbb{R}^N)$ verifica l'equazione (1.13) in ogni istante $t \in [t_0, T)$ e soddisfa la condizione iniziale $\mathbf{y}(t_0) = \mathbf{y}_0$.

Naturalmente, si potrebbe analogamente definire il concetto di problema di Cauchy *all'indietro* e di tempo iniziale di una relativa soluzione in piccolo. Osserviamo anche che in (1.17) stiamo intendendo che $\mathbf{y}'(t_0)$ indichi la derivata *destra* di $\mathbf{y}$ in t_0 .

Esercizio 1.12. Dimostrare che, data una funzione $\mathbf{y} : (a, b) \rightarrow \mathbb{R}^N$ e dati $t_0 \in (a, b)$, $\mathbf{y}_0 \in \mathbb{R}^N$ tali che $(t_0, \mathbf{y}_0) \in \Omega$, $(\mathbf{y}, (a, b))$ è soluzione in piccolo di (1.16) se e solo se $(\mathbf{y}|_{[t_0, b)}, b)$ lo è del problema (1.17) e $(\mathbf{y}|_{(a, t_0]}, a)$, intendendo a come tempo iniziale, lo è del corrispondente problema all'indietro (il punto non ovvio riguarda la regolarità C^1).

Tornando alla teoria generale, cerchiamo ora delle buone condizioni su $\mathbf{f}$ che assicurino l'esistenza e, possibilmente, l'unicità della soluzione in piccolo di (1.16); naturalmente con facili modifiche si possono trattare i problemi in avanti ed all'indietro. La condizione fondamentale ha peraltro una forma abbastanza curiosa:

Definizione 1.13. Sia $(t_0, \mathbf{y}_0) \in \Omega$. Diciamo che $\mathbf{f}$ verifica la condizione $C(t_0, \mathbf{y}_0)$ qualora esistano $\varepsilon, \delta, L > 0$ tali che

$$|\mathbf{f}(t, \mathbf{y}_1) - \mathbf{f}(t, \mathbf{y}_2)| \leq L|\mathbf{y}_1 - \mathbf{y}_2| \quad \forall t \in [t_0 - \delta, t_0 + \delta], \quad \forall \mathbf{y}_1, \mathbf{y}_2 \in \overline{B}(\mathbf{y}_0, \varepsilon). \quad (1.18)$$

Dico che $\mathbf{f}$ verifica la condizione C in Ω se verifica $C(t_0, \mathbf{y}_0)$ per ogni scelta di $(t_0, \mathbf{y}_0) \in \Omega$, ove le costanti ε, δ, L possono dipendere dal punto $(t_0, \mathbf{y}_0)$.

Si noti che è parte della definizione che sia $[t_0 - \delta, t_0 + \delta] \times \overline{B}(\mathbf{y}_0, \varepsilon) \subset \Omega$. Bisogna dire che si trovano sui libri di testo diverse varianti della (1.18). La formulazione data qui sopra ci pare tuttavia essere allo stesso tempo abbastanza maneggevole e piuttosto generale. Infatti, consente di enunciare sinteticamente il risultato fondamentale della teoria:

Teorema 1.14. (Teorema di esistenza e unicità in piccolo). Sia Ω un aperto di $\mathbb{R}^{N+1}$, $\mathbf{f} \in C^0(\Omega; \mathbb{R}^N)$, $(t_0, \mathbf{y}_0) \in \Omega$ e supponiamo che $\mathbf{f}$ verifichi la condizione $C(t_0, \mathbf{y}_0)$. Allora il Problema (1.16) ammette una soluzione in piccolo $(\mathbf{y}, (a, b))$. Inoltre, tale soluzione è unica sul dominio (a, b) .

Osservazione 1.15. Con “unica sul dominio (a, b) ” intendiamo che, se $(\mathbf{z}, (a, b))$ è una soluzione in piccolo di (1.16), allora necessariamente $\mathbf{y}(t) = \mathbf{z}(t)$ per ogni $t \in (a, b)$. Invece, chiaramente, esistono altre soluzioni $(\mathbf{w}, (c, d))$ del problema, con $(c, d) \neq (a, b)$, che sono diverse dalla $(\mathbf{y}, (a, b))$ ai sensi del formalismo usato.

Concludiamo il paragrafo sviluppando la dimostrazione del Teorema, che è basata essenzialmente sul Teorema delle contrazioni. Infatti, costruiremo una riformulazione del problema (1.16) nella forma di una *equazione integrale di Volterra*, che si presta all'applicazione di un metodo di punto fisso. Si ha infatti:

Lemma 1.16. Sia $f \in C^0(\Omega; \mathbb{R}^N)$, $(t_0, \mathbf{y}_0) \in \Omega$, $(a, b) \subset \mathbb{R}$, $t_0 \in (a, b)$, $\mathbf{y} : (a, b) \rightarrow \mathbb{R}^N$. Le seguenti condizioni sono allora equivalenti:

- (i) La coppia $(\mathbf{y}, (a, b))$ è soluzione in piccolo di (1.16);
- (ii) $\mathbf{y} \in C^0((a, b); \mathbb{R}^N)$; inoltre,

$$\mathbf{y}(t) = \mathbf{y}_0 + \int_{t_0}^t \mathbf{f}(s, \mathbf{y}(s)) ds \quad \forall t \in (a, b). \quad (1.19)$$

Prova del Lemma. Viene lasciata per esercizio, poiché consiste in una serie di verifiche immediate; l'unico punto che vale la pena evidenziare è che, nel caso $\mathbf{y}$ verifichi (ii), segue immediatamente $\mathbf{y} \in C^1((a, b); \mathbb{R}^N)$; infatti, l'integranda a secondo membro di (1.19) è certamente continua; dunque, la funzione integrale è di classe C^1 grazie al Teorema [4, VIII.1.1] di derivazione della funzione integrale. ■

Prova del Teorema. Consideriamo un numero $\delta' \in (0, \delta]$, la cui scelta sarà specificata alla fine. Poniamo $X := \overline{B}(\mathbf{y}_0, \varepsilon)$, ossia la bolla chiusa in $BC^0((t_0 - \delta', t_0 + \delta'); \mathbb{R}^N)$ di centro la funzione costante $\mathbf{y}_0$ e raggio ε , rispetto alla consueta norma $\|\cdot\|_\infty$ ed ove l' ε è quello che compare nella condizione $C(t_0, \mathbf{y}_0)$ (1.18) che è tra le ipotesi. Il nostro obiettivo è quello di scegliere δ' sufficientemente piccolo affinché l'operatore

$$\mathbf{u} \mapsto \mathcal{T}\mathbf{u}, \quad \mathcal{T}\mathbf{u}(t) := \mathbf{y}_0 + \int_{t_0}^t \mathbf{f}(s, \mathbf{u}(s)) ds \quad (1.20)$$

sia una contrazione di X in sé. Se siamo in grado di fare questo, la completezza di X (che è un sottoinsieme chiuso dello spazio di Banach BC^0) e il Teorema 1.6 ci garantiscono che la mappa $\mathcal{T}$ ammette uno ed un solo punto fisso $\mathbf{y}$. Grazie al precedente Lemma, $(\mathbf{y}, (t_0 - \delta', t_0 + \delta'))$ sarà soluzione del problema di Cauchy.

Per attuare questo programma, dobbiamo verificare due cose: che l'immagine di $\mathcal{T}$ è contenuta in X e che $\mathcal{T}$ è una contrazione. Cominciamo dalla prima: poiché $[t_0 - \delta, t_0 + \delta] \times \overline{B}(\mathbf{y}_0, \varepsilon)$ è compatto, la funzione $(t, \mathbf{y}) \mapsto |\mathbf{f}(t, \mathbf{y})|$ ammette un valore massimo, che battezziamo M , su tale insieme. Inoltre, si vede subito che $\mathcal{T}\mathbf{u}$ è continua; infine,

$\forall t \in (t_0 - \delta', t_0 + \delta')$, si ha:

$$\begin{aligned} |\mathcal{T}\mathbf{u}(t) - \mathbf{y}_0| &= \left| \int_{t_0}^t \mathbf{f}(s, \mathbf{u}(s)) ds \right| \\ &\leq \left| \int_{t_0}^t |\mathbf{f}(s, \mathbf{u}(s))| ds \right| \\ &\leq M\delta'; \end{aligned} \quad (1.21)$$

dunque, siamo sicuri che $\mathcal{T}\mathbf{u} \in X$, a patto di scegliere $\delta' \leq \varepsilon/M$.

Verifichiamo ora che $\mathcal{T}$ contrae. Siano $\mathbf{u}, \mathbf{v} \in X$. Per ogni $t \in (t_0 - \delta', t_0 + \delta')$, si ha

$$\begin{aligned} |\mathcal{T}\mathbf{u}(t) - \mathcal{T}\mathbf{v}(t)| &\leq \left| \int_{t_0}^t |\mathbf{f}(s, \mathbf{u}(s)) - \mathbf{f}(s, \mathbf{v}(s))| ds \right| \\ &\leq L \left| \int_{t_0}^t |\mathbf{u}(s) - \mathbf{v}(s)| ds \right|, \end{aligned} \quad (1.22)$$

da cui, prendendo l'estremo superiore per $s \in (t_0 - \delta', t_0 + \delta')$ della funzione nell'ultimo integrale, otteniamo

$$|\mathcal{T}\mathbf{u}(t) - \mathcal{T}\mathbf{v}(t)| \leq L \left| \int_{t_0}^t \|\mathbf{u} - \mathbf{v}\|_{\infty} ds \right| \leq L\delta' \|\mathbf{u} - \mathbf{v}\|_{\infty}, \quad \forall t. \quad (1.23)$$

A questo punto, la tesi segue prendendo anche a primo membro l'estremo superiore rispetto a t , purché imponiamo l'ulteriore vincolo $\delta' < 1/L$. ■

Val la pena osservare che un'ampia classe di funzioni $\mathbf{f}$ soddisfano la condizione C in Ω ; ad esempio tutte le funzioni di classe C^1 o, addirittura, tutte le funzioni continue, differenziabili rispetto alle variabili $\mathbf{y}$ e con relative derivate parziali continue nel complesso delle variabili $(t, \mathbf{y})$ (verificare per esercizio!). Sotto tali ipotesi abbiamo dunque che il problema di Cauchy ammette soluzione in piccolo per ogni scelta del dato $(t_0, \mathbf{y}_0)$ in Ω . Si noti inoltre che è consentito ai coefficienti ε, δ, L di dipendere da $(t_0, \mathbf{y}_0)$; dunque, ai fini dell'esistenza in piccolo, *non dà fastidio* che f o qualche sua derivata possa esplodere se ci avviciniamo al bordo di Ω .

Grazie anche a questa osservazione, utilizzando una tecnica di “bootstrap” analoga a quella vista nella dimostrazione del Lemma 1.16, si ottiene facilmente un risultato di regolarità per la soluzione in piccolo del Problema di Cauchy (1.16), ove l'unicità è ancora da intendersi nel senso reso preciso nell'Oss. 1.15:

Proposizione 1.17. *Sia $\mathbf{f}$ di classe $C^k(\Omega, \mathbb{R}^N)$, $k \geq 1$. Allora, per ogni $(t_0, \mathbf{y}_0) \in \Omega$, il Problema (1.16) ammette una ed una sola soluzione in piccolo di classe C^{k+1} .*

1.3 Prolungamento ed esistenza in grande

Abbiamo visto che la Definizione 1.9 di soluzione in piccolo non è molto maneggevole, perché costringe a “portarsi dietro” il dominio (a, b) . Per evitare questo fastidio (ma non è questo l'unico motivo), introduciamo ora la definizione di *soluzione massimale*

di un'EDO, riferendoci al caso del Problema di Cauchy in avanti (1.17) perché la trattazione risulta più agile. Il primo passo è quello di definire una relazione d'ordine nella famiglia delle soluzioni in piccolo. Per semplicità, in tutto il paragrafo, supporremo sempre, senza ripeterlo, che $\mathbf{f}$ sia continua e verifichi la condizione C in tutto Ω .

Definizione 1.18. Siano $(\mathbf{y}_1, T_1)$ e $(\mathbf{y}_2, T_2)$ due soluzioni in piccolo di (1.17). Diciamo che $(\mathbf{y}_2, T_2)$ è un prolungamento di $(\mathbf{y}_1, T_1)$ e scriviamo $(\mathbf{y}_1, T_1) \preceq (\mathbf{y}_2, T_2)$ quando $T_1 \leq T_2$ e $\mathbf{y}_1(t) = \mathbf{y}_2(t)$ per ogni $t \in [t_0, T_1]$.

Il passo fondamentale è costituito dal seguente

Lemma 1.19. La relazione $\preceq$ è di ordine totale nella famiglia delle soluzioni in piccolo.

Prova. Date due soluzioni $(\mathbf{y}_1, T_1)$ e $(\mathbf{y}_2, T_2)$, dobbiamo mostrare che esse sono confrontabili. Procediamo per assurdo e supponiamo che non lo siano. Ponendo allora $T := \min\{T_1, T_2\}$, definiamo

$$t^* := \sup \{t \in (t_0, T) \mid (\mathbf{y}_1)|_{[t_0, t]} \equiv (\mathbf{y}_2)|_{[t_0, t]}\}. \quad (1.24)$$

Grazie al Teorema 1.14, l'insieme di cui si fa il sup è non vuoto e si ha $t^* > t_0$. Osserviamo anche che deve essere $t^* < T$; altrimenti una delle due soluzioni sarebbe un prolungamento dell'altra. Dunque, t^* appartiene al dominio di entrambe le funzioni e, grazie alla continuità di queste, si ha che $\mathbf{y}_1(t^*) = \mathbf{y}_2(t^*)$. Pongo allora $\mathbf{z}_0 := \mathbf{y}_1(t^*)$. Poiché $\mathbf{f}$ verifica la condizione C , siamo certi che il problema di Cauchy in avanti

$$\begin{cases} \mathbf{y}'(t) = \mathbf{f}(t, \mathbf{y}(t)) \\ \mathbf{y}(t^*) = \mathbf{z}_0 \end{cases} \quad (1.25)$$

ammette un'unica (nel senso solito) soluzione in piccolo $(\mathbf{y}^*, t^{**})$, ove $t^{**} > t^*$. In particolare, per la proprietà di unicità, sia $\mathbf{y}_1$ che $\mathbf{y}_2$ devono coincidere con $\mathbf{y}^*$, e quindi tra di loro, in un intorno destro di t^* , il che va contro la definizione (1.24) di t^* . ■

È evidente che nel caso si consideri il problema “bilaterale” (1.16), la relazione $\preceq$ non è più di ordine totale: bisogna considerare separatamente il problema in avanti e quello all'indietro (vedi anche l'Esercizio 1.12).

Il Lemma precedente giustifica una Definizione e un Teorema:

Definizione 1.20. Una soluzione in piccolo di (1.16) (o di (1.17)) si dice *massimale* se non ammette prolungamenti propri (ossia diversi da lei stessa).

Teorema 1.21. Il Problema di Cauchy (1.16) (ovvero (1.17)) ammette una e una sola soluzione in piccolo massimale.

Prova. La dimostrazione è costruttiva e, a questo punto, immediata. Nel caso ad esempio del problema in avanti, basta definire

$$T_{\max} := \sup \{T : \exists (\mathbf{y}, T) \text{ soluzione in piccolo di (1.17)}\} \quad (1.26)$$

e, per ogni $t < T_{\max}$, porre

$$\mathbf{y}_{\max}(t) := \mathbf{y}(t), \quad (1.27)$$

ove $\mathbf{y}$ è una qualsiasi funzione che risolve (1.17) in un intervallo contenente t . In effetti una tale $\mathbf{y}$ esiste grazie alla scelta di $T_{\max}$ ed il valore di $\mathbf{y}_{\max}(t)$ è indipendente dalla scelta della specifica $\mathbf{y}$ in virtù del Lemma precedente. ■

Dunque, grazie a questo Teorema, se $\mathbf{f}$ verifica C in tutto Ω , tra le soluzioni in piccolo ne possiamo individuare una che è privilegiata in quanto non ammette prolungamenti. In questo caso, siamo autorizzati a parlare semplicemente “della” soluzione di un problema di Cauchy (1.16) oppure (1.17), intendendo che ci riferiamo a quella massimale. Molto spesso (ma non sempre), parlando della soluzione massimale, ometteremo di indicarne il dominio e scriveremo semplicemente $\mathbf{y}$, anziché $(\mathbf{y}, (a, b))$.

A questo punto, ci si potrebbe chiedere se, una volta che consideriamo la soluzione massimale di un problema di Cauchy, questa non debba magari avere un dominio “predeterminato”, ossia leggibile già dai dati $\mathbf{f}, t_0, \mathbf{y}_0$. In generale, però, questo non è vero, come dimostrano diversi esempi elementari e, d’altro canto, è anche intuibile da considerazioni geometriche (si pensi a situazioni in cui Ω ha una forma molto irregolare). Possiamo anzi dire che la possibilità di determinare a priori (cioè senza risolvere l’equazione) l’intervallo di definizione della soluzione massimale è uno dei problemi più interessanti della teoria. Esiste peraltro un criterio abbastanza comodo che permette di stabilire che una data soluzione *non* è massimale. Lo enunciamo nel caso del problema in avanti (1.17).

Proposizione 1.22. *Sia K un sottoinsieme chiuso e limitato di Ω e sia $(\mathbf{y}, (a, b))$ una soluzione in piccolo di (1.16) tale che $\text{graf } \mathbf{y} \subset K$. Allora $(\mathbf{y}, (a, b))$ non è massimale; anzi, ammette un prolungamento sia a destra di b che a sinistra di a .*

Prova. Poniamo $M := \max\{|\mathbf{f}(t, \mathbf{y})|, (t, \mathbf{y}) \in K\}$, che esiste per il Teorema di Weierstrass. Dati $r, s \in (a, b)$, si ha

$$|\mathbf{y}(s) - \mathbf{y}(r)| \leq \int_r^s |\mathbf{f}(\sigma, \mathbf{y}(\sigma))| d\sigma \leq M|s - r|; \quad (1.28)$$

dunque, $\mathbf{y}$ è Lipschitz in (a, b) e, grazie alla Prop. 1.7, ammette un prolungamento Lipschitz definito su $[a, b]$. Inoltre, $(b, \mathbf{y}(b)) \in K$, poiché K è chiuso. Risolvendo il problema di Cauchy con dato $(b, \mathbf{y}(b)) \in K \subset \Omega$, si riesce allora a prolungare $\mathbf{y}$ a destra di b . Allo stesso modo si procede a sinistra di a . ■

Un altro modo per stabilire se una soluzione è massimale è connesso col concetto di *soluzione in grande*, che ora introduciamo. Presentiamo anzi subito un nuovo Teorema di esistenza e unicità. Come è lecito aspettarsi in base alle considerazioni precedenti, entrano qui in gioco due nuove ipotesi, che coinvolgono rispettivamente la geometria del dominio Ω e la regolarità di $\mathbf{f}$.

Teorema 1.23. (Teorema di esistenza e unicità in grande). *Supponiamo $\Omega = (a, b) \times \mathbb{R}^N$ e che $\mathbf{f} \in C^0(\Omega, \mathbb{R}^N)$ verifichi C in Ω . Supponiamo inoltre che esistano due funzioni $\alpha, \beta : (a, b) \rightarrow \mathbb{R}$ continue, non negative e tali che*

$$|\mathbf{f}(t, \mathbf{y})| \leq \alpha(t)|\mathbf{y}| + \beta(t) \quad \forall t \in (a, b), \quad \forall \mathbf{y} \in \mathbb{R}^N. \quad (1.29)$$

Infine, sia $t_0 \in (a, b)$, $\mathbf{y}_0 \in \mathbb{R}^N$. Allora il Problema (1.16) ammette una ed una sola soluzione massimale $\mathbf{y}$ definita su (a, b) .

Osservazione 1.24. Notiamo che la continuità delle funzioni α e β richiesta nella (1.29) è puramente di comodo; è sufficiente che tali funzioni siano *localmente limitate*, ossia limitate su ogni sottoinsieme chiuso e limitato di (a, b) . Ancora una volta non dà fastidio che possano “sparare” agli estremi. Osserviamo anche che la (1.29) è detta a volte condizione di *sottolinearità*: in effetti, vogliamo che a t fissato la funzione $\mathbf{y} \rightarrow |f(t, \mathbf{y})|$ abbia crescita al più lineare.

Dunque una soluzione $\mathbf{y}$ si dice “in grande” quando Ω è una striscia verticale $(a, b) \times \mathbb{R}^N$ (altrimenti il concetto non ha senso) ed il dominio di $\mathbf{y}$ è tutto (a, b) . Esempi visti in precedenza mostrano che, anche nel caso $\Omega = \mathbb{R}^{N+1}$, la condizione C da sola non garantisce che la soluzione massimale di (1.16) sia definita in grande.

Premettiamo alla prova del teorema un risultato che ci sarà utile anche in altre occasioni, notando che un enunciato analogo vale anche per il problema all’indietro.

Lemma 1.25. (Lemma di Gronwall in avanti). Siano $\gamma \geq 0$ e siano date due funzioni $m, \varphi \in C^0([t_0, b))$ tali che $m \geq 0$. Sia inoltre

$$\varphi(t) \leq \gamma + \int_{t_0}^t m(s)\varphi(s) ds \quad \forall t \in [t_0, b). \quad (1.30)$$

Allora si ha che

$$\varphi(t) \leq \gamma \exp\left(\int_{t_0}^t m(s) ds\right) \quad \forall t \in [t_0, b). \quad (1.31)$$

Prova del Lemma. Pongo

$$M(t) := \int_{t_0}^t m(s) ds; \quad G(t) := \exp(-M(t)) \int_{t_0}^t m(s)\varphi(s) ds.$$

Procedendo con un calcolo diretto è allora facile verificare che per ogni tempo t vale la disuguaglianza

$$G'(t) \leq \gamma m(t) \exp(-M(t)),$$

da cui, essendo $G(t_0) = 0$, integrando segue subito

$$G(t) \leq \gamma(1 - \exp(-M(t))).$$

A questo punto la tesi segue facilmente moltiplicando per $\exp(M(t))$ e utilizzando ancora l’ipotesi (1.30). ■

Prova del Teorema 1.23. Supponiamo per assurdo che il *tempo finale* della soluzione massimale di (1.16) sia $\kappa < b$. Dall’equazione integrale (1.19) ho subito

$$\begin{aligned} |\mathbf{y}(t)| &\leq |\mathbf{y}_0| + \int_{t_0}^t (\alpha(s)|\mathbf{y}(s)| + \beta(s)) ds \\ &\leq |\mathbf{y}_0| + (\kappa - t_0) \max_{[t_0, \kappa]} \beta + \int_{t_0}^t \alpha(s)|\mathbf{y}(s)| ds, \quad \forall t \in [t_0, \kappa) \end{aligned} \quad (1.32)$$

da cui, applicando Gronwall con le scelte

$$\varphi = |\mathbf{y}|, \quad \gamma = |\mathbf{y}_0| + (\kappa - t_0) \max_{[t_0, \kappa]} \beta, \quad m = \alpha,$$

segue subito

$$|\mathbf{y}(t)| \leq (|\mathbf{y}_0| + (\kappa - t_0) \max_{[t_0, \kappa]} \beta) \exp \left(\int_{t_0}^{\kappa} \alpha(s) ds \right) \quad \forall t \in [t_0, \kappa]; \quad (1.33)$$

conseguentemente esiste $M > 0$ tale che

$$\text{graf } \mathbf{y} \subset [t_0, \kappa] \times \overline{B}(t_0, M),$$

che è evidentemente un compatto contenuto in Ω . La Prop. 1.22 assicura allora che $\mathbf{y}$ non è massimale, il che fornisce l'assurdo. ■

Vale la pena concludere questa sezione ricordando alcune condizioni semplici da verificare che garantiscono l'applicabilità del Teorema 1.23 (ovviamente siamo sempre nel caso in cui Ω è una striscia verticale). Ad esempio, le ipotesi sono soddisfatte se $\mathbf{f}$ è continua ed esiste $L \in C^0((a, b); \mathbb{R}^+)$ tale che

$$|\mathbf{f}(t, \mathbf{y}_1) - \mathbf{f}(t, \mathbf{y}_2)| \leq L(t) |\mathbf{y}_1 - \mathbf{y}_2| \quad \forall t \in (a, b), \quad \mathbf{y}_1, \mathbf{y}_2 \in \mathbb{R}^N; \quad (1.34)$$

infatti è semplice controllare che tale condizione implica sia la (1.29) sia la condizione C in tutto Ω . Anzi, come nel caso generale basta che L sia localmente limitata. Ribadiamo infine che, la regolarità $\mathbf{f} \in C^1(\Omega; \mathbb{R}^N)$ da sola non è sufficiente a garantire la (1.29) (e dunque l'esistenza di una soluzione definita in grande).

1.4 Dipendenza continua dai dati

Ci chiediamo ora come varia il comportamento di una successione di problemi di Cauchy al variare dei dati. L'importanza di una tale questione è chiara anche da un punto di vista pratico: se un'equazione descrive un qualche fenomeno fisico, può esserci incertezza nella misurazione precisa del dato iniziale, oppure nella determinazione dell'espressione della $\mathbf{f}$. Dunque, in questo caso è importante che a “piccole” variazioni dei dati corrispondano “piccole” variazioni della soluzione; altrimenti, il sistema che stiamo considerando rischia di essere un pessimo modello della situazione fisica.

In realtà, sotto le ipotesi del Teorema di esistenza in piccolo le cose si comportano abbastanza bene, anche se il problema è più complesso di come potrebbe parere a prima vista, e ora vediamo il perché. Consideriamo, al variare di $k \in \mathbb{N}$, la famiglia di problemi

$$\begin{cases} \mathbf{y}'_k(t) = \mathbf{f}_k(t, \mathbf{y}_k(t)) \\ \mathbf{y}_k(t_0) = \mathbf{y}_{0,k}, \end{cases} \quad (1.35)$$

ove supponiamo che siano date:

- (a) una successione $\{\mathbf{f}_k\}$ di funzioni almeno della regolarità $\mathbf{f}_k \in C^0(\Omega; \mathbb{R}^N)$ (dunque, il dominio Ω è, per semplicità, lo stesso per tutte le $\mathbf{f}_k$);
- (b) una corrispondente famiglia di dati iniziali $\{\mathbf{y}_{0,k}\} \subset \mathbb{R}^N$, col vincolo $(t_0, \mathbf{y}_{0,k}) \in \Omega$ per ogni $k \in \mathbb{N}$. In particolare, lasciamo fisso il punto t_0 ;
- (c) opportuni dati limite $\mathbf{f} \in C^0(\Omega; \mathbb{R}^N)$, $(t_0, \mathbf{y}_0) \in \Omega$, ove supponiamo già da subito che sia $\mathbf{y}_{0,k} \rightarrow \mathbf{y}_0$ (nel senso della convergenza usuale in $\mathbb{R}^N$).

Per poter applicare il Teorema di esistenza e unicità in piccolo, inoltre, supponiamo:

(d) le $\mathbf{f}_k$ verifichino C in Ω *uniformemente rispetto a* k^1 .

Se valgono (a-d), il problema (1.35) ammette una soluzione in piccolo massimale $\mathbf{y}_k$ per ogni $k \in \mathbb{N}$.

La domanda che ci poniamo ora è duplice:

- quale tipo di convergenza $\mathbf{f}_k \rightarrow \mathbf{f}$ è opportuno supporre per avere la convergenza delle $\mathbf{y}_k$ alla soluzione (se esiste) $\mathbf{y}$ del problema limite (che ha la usuale forma (1.16))?
- corrispondentemente, in quale senso potrà valere la convergenza $\mathbf{y}_k \rightarrow \mathbf{y}$?

Prima di presentare il risultato che risponderà a tali questioni, premettiamo qualche considerazione. Notiamo innanzitutto che, sotto queste ipotesi, è sufficiente la convergenza puntuale $\mathbf{f}_k \rightarrow \mathbf{f}$ affinché il problema limite ammetta soluzione in piccolo; infatti, $\mathbf{f}$ è continua grazie alla (c) e soddisfa C in Ω come si vede passando al limite nella condizione C scritta per $\mathbf{f}_k$ (e qui è essenziale l'uniformità in k). Tuttavia, semplici esempi mostrano che, se $\mathbf{f}_k \rightarrow \mathbf{f}$ solo puntualmente, non è affatto detto che $\mathbf{y}_k \rightarrow \mathbf{y}$, neanche in senso puntuale; anzi, il limite di $\mathbf{y}_k$ potrebbe essere una funzione discontinua o addirittura non esistere.

D'altro canto, la convergenza uniforme su Ω sembra una condizione troppo forte, come mostra un semplicissimo

Esempio 1.26. Siano $\Omega = \mathbb{R}^2$, $f_k(t, y) = 1/k$ (cosicché $f = 0$ e $f_k \rightarrow 0$ uniformemente in $\mathbb{R}^2$), $t_0 = y_{0,k} = y_0 = 0$. Allora, ovviamente, si ha che $y_k(t) = t/k$ e $y(t) = 0$, con dominio $\mathbb{R}$ per entrambe. D'altronde y_k non tende a 0 uniformemente in $\mathbb{R}$.

In effetti, la questione della corretta nozione di convergenza per $\mathbf{f}_k$ si riconduce al discorso del § 1.1 sulla struttura dello spazio delle funzioni continue definite su un aperto Ω . In quell'occasione avevamo imposto a priori una proprietà di limitatezza introducendo lo spazio $BC^0(\Omega; \mathbb{R}^N)$; a quel punto, la norma $\|\cdot\|_\infty$ forniva la nozione più naturale di convergenza, che rendeva BC^0 uno spazio completo. Ora, tuttavia, non c'è motivo di chiedere che le $\mathbf{f}_k$ siano limitate: resterebbero fuori molti casi importanti. Serve un approccio diverso, che porta a introdurre una nuova nozione di convergenza:

Definizione 1.27. Dico che la successione $\{\mathbf{f}_k\}$ converge a $\mathbf{f}$ sui compatti di Ω (ovvero che converge in $C^0(\Omega; \mathbb{R}^N)$) se per ogni sottoinsieme compatto $K \subset \Omega$ si ha che $(\mathbf{f}_k)|_K \rightarrow \mathbf{f}|_K$ uniformemente.

Si noti che la nozione sopra introdotta è piuttosto naturale per lo spazio $C^0(\Omega; \mathbb{R}^N)$: infatti conserva la continuità (verificare!) e permette che le funzioni $\mathbf{f}_k, \mathbf{f}$ possano essere non limitate; tuttavia non esiste alcuna norma che induce su $C^0(\Omega; \mathbb{R}^N)$ tale convergenza². Tornando all'Esempio 1.26, possiamo vedere che in effetti le soluzioni $y_k = t/k$ convergono a 0 in $C^0(\mathbb{R})$. Enunciamo allora il risultato preciso, la cui formulazione è complicata dal fatto che vogliamo considerare soluzioni in piccolo:

¹ossia ammettiamo che i numeri ε, δ, L che compaiono nella C possano dipendere dal punto di Ω , ma non da k

²per chi conosce la teoria degli spazi metrici, notiamo che la convergenza sui compatti è indotta da una distanza che rende $C^0(\Omega; \mathbb{R}^N)$ uno spazio metrico completo

Teorema 1.28. *Sia data la famiglia di problemi di Cauchy (1.35), ove i dati $\mathbf{f}_k$, t_0 , $\mathbf{y}_{0,k}$ verificano le condizioni (a), (b), (d) sopra riportate. Supponiamo inoltre che esistano $\mathbf{f} : \Omega \rightarrow \mathbb{R}^N$, $\mathbf{y}_0 \in \mathbb{R}^N$ con $(t_0, \mathbf{y}_0) \in \Omega$ tali che*

$$\mathbf{f}_k \rightarrow \mathbf{f} \quad \text{in } C^0(\Omega; \mathbb{R}^N), \quad \mathbf{y}_{0,k} \rightarrow \mathbf{y}_0. \quad (1.36)$$

Dette allora $\mathbf{y}_k$, $\mathbf{y}$ le soluzioni massimali dei problemi di Cauchy (1.35) e (1.16), si ha che per ogni intervallo chiuso e limitato $I \subset \text{dom } \mathbf{y}$ contenente t_0 al suo interno esiste $\bar{k}$ tale che $\forall k \geq \bar{k}$ si ha che $I \subset \text{dom } \mathbf{y}_k$; inoltre, $\mathbf{y}_k$ tende a $\mathbf{y}$ uniformemente su I .

Si osservi che la continuità di $\mathbf{f}$ e la condizione C per la $\mathbf{f}$ stessa sono ora conseguenza delle ipotesi (a), (d), (1.36); dunque, non serve supporre la (c). Inoltre, anche se il significato della tesi è grosso modo che “ $\mathbf{y}_k$ tende a $\mathbf{y}$ sui compatti”, abbiamo dovuto essere un po’ più precisi nell’enunciato perché il dominio di $\mathbf{y}_k$ può variare (anche se “non troppo”) con k . La dimostrazione del teorema è interessante, anche se un po’ laboriosa:

Prova. Mettiamoci nel caso $N = 1$ che permette di evitare alcune complicazioni puramente tecniche. Usiamo dunque le notazioni scalari f_k, y_k , ecc.. Sia dato un nuovo intervallo chiuso e limitato J tale che $I \subset \text{int } J$, $J \subset \text{dom } y$. Poiché $\text{graf } f|_J$ è un compatto contenuto in Ω , è facile mostrare che esiste $\varepsilon > 0$ tale che l’insieme

$$K_\varepsilon := \{(t, x) \in \mathbb{R}^2 : t \in J, x \in [y(t) - \varepsilon, y(t) + \varepsilon]\} \quad (1.37)$$

è contenuto in Ω . Ponendo $\Lambda := \text{int } J \times \mathbb{R}$, definiamo, per $(t, x) \in \bar{\Lambda} = J \times \mathbb{R}$,

$$\bar{f}_k(t, x) := \begin{cases} f_k(t, x) & \text{se } (t, x) \in K_\varepsilon, \\ f_k(t, y(t) - \varepsilon) & \text{se } x < y(t) - \varepsilon, \\ f_k(t, y(t) + \varepsilon) & \text{se } x > y(t) + \varepsilon, \end{cases} \quad (1.38)$$

e analogamente costruiamo $\bar{f}$ a partire da f . È chiaro allora che $(y, \text{int } J)$ è soluzione in piccolo del problema di Cauchy (1.16) ove f è rimpiazzata da $\bar{f}$. Le due funzioni coincidono infatti su un intorno del grafico di y . Invece, a priori y_k non è legata a $\bar{f}_k$ ed infatti ci tocca definire una nuova famiglia di funzioni. Chiamiamo $\bar{y}_k$ la soluzione del nuovo problema di Cauchy

$$\begin{cases} \bar{y}'_k(t) = \bar{f}_k(t, \bar{y}_k(t)) \\ \bar{y}_k(t_0) = y_{0,k}, \end{cases} \quad (1.39)$$

ove guardiamo $\bar{f}_k$ come funzione da Λ in $\mathbb{R}$. In effetti, siamo sicuri che (1.39) ha soluzione in grande; infatti le $\bar{f}_k$ sono tutte continue, limitate e Lipschitziane in y con costante di Lipschitz L indipendente da t e da k . La verifica di questo fatto, che omettiamo, è facile ma tecnica e si basa sulla condizione (d) e sul fatto che con la troncatura ci siamo essenzialmente ridotti a guardare le $\bar{f}_k$ sul compatto K_ε .

A priori non c’è motivo per cui debba essere $y_k = \bar{y}_k$. Pazienza: cominciamo lo stesso a dimostrare che $\bar{y}_k \rightarrow y$ uniformemente in I . Utilizzando l’equazione integrale

di Volterra, si vede che per $t \in I$ è

$$\begin{aligned} |\bar{y}_k(t) - y(t)| &\leq |y_{0,k} - y_0| + \int_{t_0}^t |\bar{f}_k(s, \bar{y}_k(s)) - f(s, y(s))| ds \\ &\leq |y_{0,k} - y_0| + \int_{t_0}^t |\bar{f}_k(s, \bar{y}_k(s)) - \bar{f}_k(s, y(s))| ds \\ &\quad + \int_{t_0}^t |\bar{f}_k(s, y(s)) - f(s, y(s))| ds \end{aligned} \quad (1.40)$$

e vogliamo ora stimare i due integrali a secondo membro. Per quanto riguarda il primo, si ha

$$\int_{t_0}^t |\bar{f}_k(s, \bar{y}_k(s)) - \bar{f}_k(s, y(s))| ds \leq L \int_{t_0}^t |\bar{y}_k(s) - y(s)| ds. \quad (1.41)$$

Per quanto riguarda il secondo integrale, notiamo che $\bar{f}_k$ coincide con f_k su K_ε e in particolare su graf $y|_I$, che è compatto. Inoltre, sappiamo che $\bar{f}_k$ tende a $\bar{f}$ sui compatti di Λ ; e dunque uniformemente su graf $y|_I$. Segue che esiste una successione $\{\varepsilon_k\}$ non negativa, infinitesima e tale che

$$\int_{t_0}^t |\bar{f}_k(s, y(s)) - f(s, y(s))| ds \leq \varepsilon_k \quad \forall t \in I. \quad (1.42)$$

Tenendo conto di (1.41–1.42), si vede che si può applicare il lemma di Gronwall nella formula (1.40), ottenendo

$$|\bar{y}_k(t) - y(t)| \leq (|y_{0,k} - y_0| + \varepsilon_k) \exp(L \text{lung}(I)) =: \delta_k \quad \forall t \in I, \quad (1.43)$$

ove la successione $\{\delta_k\}$ è chiaramente infinitesima. Questo mostra la convergenza uniforme $\bar{y}_k \rightarrow y$ in I .

Per finire, chiamando $\bar{k}$ il minimo indice tale che $\delta_k < \varepsilon$ per ogni $k \geq \bar{k}$ (e per cui si ha, in particolare, $|y_{0,k} - y_0| < \varepsilon$), si ha che per ogni $k \geq \bar{k}$ è $y_k \equiv \bar{y}_k$ in I ; infatti, gli elementi della successione $\{\bar{y}_k\}$ sono vincolati ad avere grafici contenuti in K_ε , dove f_k e $\bar{f}_k$ coincidono. Quindi, per tali k , y_k e $\bar{y}_k$ sono soluzioni dello stesso problema di Cauchy, e dunque anch'esse coincidono. ■

1.5 Alcuni tipi di equazioni

Sviluppiamo in questo paragrafo sia nozioni teoriche sia strumenti pratici di risoluzione che si applicano ad alcuni tipi particolari di equazioni. Il primo che presentiamo è indubbiamente anche il più importante:

Equazioni a variabili separabili. Hanno questo nome le equazioni *scalari* della forma

$$y'(t) = a(t)b(y(t)), \quad (1.44)$$

ove siano dati I, J intervalli *aperti*, non necessariamente limitati, di $\mathbb{R}$ e due funzioni $a \in C^0(I; \mathbb{R})$, $b \in C^0(J; \mathbb{R})$. Il motivo per cui una tale equazione si chiama a variabili separabili è chiaro: dividendo entrambi i membri per $b(y(t))$ è formalmente possibile sviluppare la relazione ottenuta integrando in tempo.

Un tale procedimento formale necessita tuttavia di essere reso rigoroso alla luce della teoria descritta finora. Innanzitutto, osserviamo che, ponendo in modo naturale $\Omega := I \times J$ e $f(t, y) := a(t)b(y)$, la f risulta senz'altro continua; inoltre, se b è *localmente Lipschitziana* in J , la condizione C risulta verificata in tutto Ω ed è possibile applicare il Teorema di esistenza ed unicità in piccolo. Se inoltre $J = \mathbb{R}$ e vale l'ipotesi di *sottolinearità*, ossia esistono $C_1, C_2 > 0$ tali che

$$|b(y)| \leq C_1 + C_2|y| \quad \forall y \in \mathbb{R}, \quad (1.45)$$

allora la soluzione massimale y risulta definita in grande e cioè in tutto I . Infatti, (1.45) implica chiaramente (1.29) con ovvie scelte di α e β .

Tuttavia, la tecnica di risoluzione che descriveremo in dettaglio tra un attimo consentirà di calcolare esplicitamente il dominio di y una volta che sia fissato il dato iniziale $(t_0, y_0) \in I \times J$. A tale proposito, osserviamo innanzitutto che, qualora $b(y_0) = 0$, allora la soluzione del problema di Cauchy per (1.44) risulta definita su tutto I e costantemente uguale a y_0 . In caso contrario, supponendo che y sia una soluzione, notiamo che non può essere $b(y(t)) = 0$ per alcun tempo t ; altrimenti, il grafico di y verrebbe a intersecare il grafico di una soluzione costante. Dunque, ha senso scrivere

$$a(t) = \frac{y'(t)}{b(y(t))}, \quad (1.46)$$

dato che il denominatore a secondo membro non si può mai annullare. In un certo senso, una volta che siano fissati t_0 ed y_0 , ci interessa soltanto il comportamento di b nel più grande sottointervallo (necessariamente aperto) $J_0 \subset J$ contenente y_0 in cui b non si annulla. In effetti, abbiamo la seguente

Proposizione 1.29. *Siano $a \in C^0(I; \mathbb{R})$, $b \in \text{Lip}_{\text{loc}}(J; \mathbb{R})$. Sia $(t_0, y_0) \in I \times J$ e sia $b(y_0) \neq 0$. Posto per $t \in I$, $y \in J_0$*

$$A(t) := \int_{t_0}^t a(s) ds, \quad B(y) := \int_{y_0}^y \frac{du}{b(u)}, \quad (1.47)$$

allora la soluzione massimale del problema di Cauchy per (1.44) con la scelta del dato (t_0, y_0) è data da $(B^{-1}(A(t)), I_0)$, ove I_0 è il più grande sottointervallo aperto di I contenente t_0 e tale che $A(I_0) \subset B(J_0)$.

Prova. Poichè b non cambia segno in J_0 , la funzione B è strettamente monotona, e dunque invertibile, in J_0 . Ne consegue che la funzione $t \mapsto y(t) := B^{-1}(A(t))$ è definita su I_0 . Inoltre è di classe C^1 poichè lo sono A e B^{-1} . Due verifiche dirette consentono di mostrare che y risolve l'equazione e soddisfa la condizione iniziale.

Infine, y è massimale. Per vederlo, mostriamo che prolunga ogni altra soluzione (z, I') del problema di Cauchy. In effetti, dato $t \in I'$, ad esempio con $t > t_0$, non potendo $b(z)$ mai annullarsi in $[t_0, t]$, l'equazione mi dice che

$$A(t) = \int_{t_0}^t a(s) ds = \int_{t_0}^t \frac{ds}{b(z(s))} = B(z(t)),$$

da cui subito $z(t) = B^{-1}(A(t))$ e dunque $t \in I_0$ e $z(t) = y(t)$. ■

Osservazione 1.30. Si noti che il risultato dimostrato qui sopra permette di calcolare esplicitamente la soluzione del Problema di Cauchy per (1.44), purché si sappiano integrare le funzioni a e b e si riesca a calcolare l'inversa B^{-1} , il che naturalmente può essere impresa ancora più ardua del calcolo dei due integrali.

Equazioni omogenee. Si tratta di equazioni *scalari* della forma

$$y'(t) = f(y(t)/t). \quad (1.48)$$

È possibile trovarne una soluzione tramite la sostituzione, abbastanza intuitiva, $v := y/t$. In questo modo, la (1.48) si trasforma in

$$v'(t) = \frac{f(v(t))v(t)}{t}, \quad (1.49)$$

che è un'equazione a variabili separabili. A questo punto, se si deve studiare un problema di Cauchy per la (1.48), in genere la cosa più comoda è trasformare anche il dato iniziale ponendo naturalmente $v(t_0) := y(t_0)/t_0$; il denominatore non può annullarsi perchè l'equazione (1.48) non ha significato per $t = 0$.

Osservazione 1.31. Si noti che sono del tipo (1.48) tutte le equazioni della forma

$$y'(t) = \frac{P(t, y)}{Q(t, y)}, \quad (1.50)$$

ove P e Q sono polinomi dello stesso grado k *omogenei* (ossia tutti i termini hanno grado k); basta infatti dividere numeratore e denominatore per t^k . Osserviamo peraltro che in questo caso, se il denominatore di (1.50) non è divisibile per t , è lecito imporre una c.i. della forma $y(0) = y_0$ con $y_0 \neq 0$. Dunque, può capitare che esistano soluzioni della (1.50) definite anche in (un intorno di) 0.

Esercizio 1.32. Studiare qualitativamente le soluzioni dell'equazione

$$y'(t) = \frac{t + y}{t - y} \quad (1.51)$$

al variare della condizione iniziale $y(t_0) = y_0$.

Equazioni di ordine superiore al primo. Finora abbiamo sviluppato la teoria solo per le equazioni del primo ordine. Dovrebbe peraltro essere ben noto (vedi [4, § 9.3]) che è possibile affrontare lo studio dell'equazione scalare di ordine $k > 1$ in forma normale

$$y^{(k)}(t) = f(t, y(t), \dots, y^{(k-1)}(t)) \quad (1.52)$$

attraverso la sostituzione

$$y_1 := y, \quad y_2 := y', \quad \dots, \quad y_k := y^{(k-1)}, \quad (1.53)$$

che permette di trasformare la (1.52) in un sistema del primo ordine del tipo (1.13) per la variabile vettoriale $\mathbf{y} := (y_1, \dots, y_k)$, ove naturalmente si è posto:

$$\mathbf{f}(t, \mathbf{y}) := (y_2, \dots, y_k, f(t, y_1, \dots, y_k)). \quad (1.54)$$

A questo punto si potrebbero scrivere vari teoremi di esistenza, unicità, regolarità, prolungamento, dipendenza continua per l'equazione (1.52) provando ad adattare la teoria vista per i sistemi del primo ordine. In realtà, si vede subito che le modifiche richieste sono di pochissimo conto e non compaiono ipotesi nuove; infatti, le prime $k - 1$ componenti della $\mathbf{f}$ “non danno fastidio”, poiché sono di classe C^∞ con crescita lineare. Dunque le ipotesi da richiedere saranno del tutto analoghe a quelle viste per i sistemi del primo ordine; pertanto omettiamo di esplicitare gli enunciati e lasciamo i dettagli per il lettore.

Equazioni autonome del secondo ordine. Vedremo nel prossimo paragrafo cosa si intende in generale per “equazione autonoma”. Per il momento, vogliamo semplicemente sviluppare una tecnica “pratica” che consenta di risolvere o quantomeno di “semplificare” le equazioni (scalari) della forma

$$y''(t) = f(y(t), y'(t)) \quad (1.55)$$

o, più in generale,

$$f(y(t), y'(t), y''(t)) = 0. \quad (1.56)$$

Si noti che la (1.56) è un'equazione *non in forma normale*. Dunque, anche nel caso f sia molto regolare potrebbero non valere i risultati teorici di esistenza e unicità delle soluzioni.

Il trucco consiste nel supporre la funzione y localmente invertibile nell'intorno di ogni punto t ; il procedimento potrà essere poi giustificato a posteriori. Proviamo allora a cercare la funzione inversa $t = t(y)$. Ponendo $v = v(y) := y'(t(y))$, si ha subito

$$y'' = \frac{dy'}{dt} = \frac{dy'}{dy} \frac{dy}{dt} = v'(y)y'(t(y)) = v'(y)v(y); \quad (1.57)$$

dunque ottengo da (1.55) la nuova equazione

$$v'(y) = \frac{f(y, v(y))}{v(y)} \quad (1.58)$$

nella variabile indipendente y (o un'espressione simile nel caso (1.56)). A questo punto, non è detto che si riesca a risolvere (1.57); tuttavia ci si è quantomeno ricondotti ad un'equazione del primo ordine. Notiamo anche che qualora sia dato un problema di Cauchy per la (1.55), sarà da imporre dapprima la condizione $v(y_0) = y_1$. In secondo luogo, una volta determinata la v , si potrà ricavare la y per integrazione imponendo la condizione $y(t_0) = y_0$. Invitiamo il lettore a svolgere esercizi su questo tipo di equazioni: in generale, quando si ha a che fare con procedimenti tecnici, la pratica aiuta più delle considerazioni teoriche.

1.6 Equazioni autonome – rappresentazione delle soluzioni

Si dicono autonome le equazioni della forma

$$\mathbf{y}'(t) = \mathbf{f}(\mathbf{y}(t)), \quad \mathbf{f} : \Omega \rightarrow \mathbb{R}^N, \quad (1.59)$$

quelle cioè in cui la $\mathbf{f}$ a secondo membro non dipende esplicitamente dal tempo. In questo paragrafo, Ω sarà quindi un aperto di $\mathbb{R}^N$ (e non di $\mathbb{R}^{N+1}$ e $f : \Omega \rightarrow \mathbb{R}^N$ almeno continua. È evidente che per tornare alle notazioni usuali, basta porre $\Lambda := \mathbb{R} \times \Omega \subset \mathbb{R}^{N+1}$, porre

$$\mathbf{g} : \Lambda \rightarrow \mathbb{R}^N, \quad \mathbf{g}(t, \mathbf{y}) := \mathbf{f}(\mathbf{y}), \quad \forall (t, \mathbf{y}) \in \Lambda, \quad (1.60)$$

e riscrivere la (1.59) nella forma equivalente $\mathbf{y}'(t) = \mathbf{g}(t, \mathbf{y}(t))$.

Per maggiore chiarezza può essere opportuno esplicitare in questo caso particolare le condizioni che assicurano l'applicabilità dei vari teoremi, scrivendole in termini di $\mathbf{f}$. Si ha che:

- $\mathbf{f}$ è Lipschitziana in un intorno di un punto $\mathbf{y}_0 \in \Omega$ se e solo se $\mathbf{g}$ verifica la condizione C in ogni punto della tipo $(t_0, \mathbf{y}_0)$ con $t_0 \in \mathbb{R}$;
- Λ è di tipo striscia se e solo se $\Omega = \mathbb{R}^N$; in questo caso, chiaramente, $\Lambda = \mathbb{R} \times \Omega = \mathbb{R}^{N+1}$ e, se c'è soluzione in grande, questa sarà definita $\forall t \in \mathbb{R}$;
- la condizione di sottolinearità (1.29) può essere riscritta in una forma (vettoriale) analoga alla (1.45).

Essenzialmente, dalle equivalenze sopra riportate, si vede che per la (1.59) le cose o funzionano per tutti i tempi o non funzionano mai. Questa uniformità del loro comportamento rispetto alla "scelta dei tempi" è ulteriormente precisata dal prossimo risultato. Si noti comunque che l'equazione (1.59) (che, del resto, nel caso scalare è un caso particolare di equazione a variabili separabili) continua a presentare tutte le patologie proprie del caso generale non lineare in termini di possibile nonunicità, dominio della soluzione non sempre determinabile a priori, ecc..

Proposizione 1.33. *Supponiamo che valgano le condizioni per l'esistenza e unicità in piccolo in tutto Ω e siano $(\mathbf{y}_1, \text{dom } \mathbf{y}_1)$, $(\mathbf{y}_2, \text{dom } \mathbf{y}_2)$ due soluzioni massimali di due distinti problemi di Cauchy per la (1.59). Supponiamo inoltre che esistano $t_1 \in \text{dom } \mathbf{y}_1$, $t_2 \in \text{dom } \mathbf{y}_2$ tali che*

$$\mathbf{y}_1(t_1) = \mathbf{y}_2(t_2). \quad (1.61)$$

Allora si ha che

$$\text{dom } \mathbf{y}_1 = \text{dom } \mathbf{y}_2 - (t_2 - t_1), \quad (1.62)$$

$$\mathbf{y}_1(t) = \mathbf{y}_2(t + t_2 - t_1) \quad \forall t \in \text{dom } \mathbf{y}_1. \quad (1.63)$$

Prova. Consideriamo la funzione $\mathbf{z}(t) := \mathbf{y}_2(t + t_2 - t_1)$, per $t \in \text{dom } \mathbf{z} := \text{dom } \mathbf{y}_2 - t_2 + t_1$. Poiché $\text{dom } \mathbf{y}_2$ è aperto e $t_2 \in \text{dom } \mathbf{y}_2$, $\mathbf{z}$ è definita almeno in un intorno I di t_1 . Si verifica facilmente che, almeno per $t \in I$,

$$\mathbf{z}'(t) = \mathbf{y}_2'(t + t_2 - t_1) = \mathbf{f}(\mathbf{y}_2(t + t_2 - t_1)) = \mathbf{f}(\mathbf{z}(t)); \quad (1.64)$$

si ha inoltre che $\mathbf{z}(t_1) = \mathbf{y}_2(t_2) = \mathbf{y}_1(t_1)$, e dunque $(\mathbf{z}, \text{dom } \mathbf{z})$ è soluzione in piccolo del problema di Cauchy per la (1.59) con la c.i. $\mathbf{z}(t_1) = \mathbf{y}_1(t_1)$. Inoltre, $\mathbf{z}$ è massimale; se infatti ammettesse un prolungamento proprio, sarebbe possibile prolungare anche $\mathbf{y}_2$, che è massimale per ipotesi. Dunque, $\mathbf{z}$ deve necessariamente coincidere con la soluzione massimale di (1.64), che già sappiamo per ipotesi essere $\mathbf{y}_1$. ■

Spieghiamo brevemente il “succo” del precedente risultato al di là dei tecnicismi nella dimostrazione. Il punto importante è che, se conosciamo una soluzione massimale $\mathbf{y}$ di un qualsiasi problema di Cauchy per la (1.59), allora, per ogni $\tau \in \mathbb{R}$, la *traslata* $t \mapsto \mathbf{y}(t + \tau)$ è ancora una soluzione della (1.59), il cui dominio si ottiene traslando (a sinistra) di τ il dominio della $\mathbf{y}$. Per questo motivo, possiamo convenzionalmente fissare a 0 l'*origine dei tempi* ponendo $t_0 = 0$; le soluzioni relative a scelte diverse del tempo iniziale si troveranno tramite traslazioni.

Vediamo ora quali sono le conseguenze geometriche di questo fatto, cominciando ad esaminare caso scalare. Qui, in generale, le soluzioni della (1.59) vengono rappresentate disegnandone il grafico nel piano (t, y) . Dunque, se conosciamo il grafico di una particolare soluzione y della (1.59), ne possiamo ottenere infiniti altri traslandolo a destra e a sinistra di un arbitrario valore $\tau \in \mathbb{R}$.

Esercizio 1.34. Sia data una soluzione y di (1.59), di cui supponiamo di conoscere il grafico. Sotto quali condizioni su y possiamo dire di conoscere il grafico di *ogni* soluzione dell'equazione (ossia per ogni scelta ammissibile del dato (t_0, y_0))?

Esercizio 1.35. È possibile che una soluzione y della (1.59) abbia punti di estremo relativo? Cosa possiamo dire del segno di y' ?

Nel caso vettoriale, non è più possibile rappresentare il grafico della generica soluzione di un'EDO. Tuttavia, almeno nel caso autonomo e se $N = 2$ (in realtà, anche con $N = 3$, a patto di essere capaci di “vedere” in 3 dimensioni), si può provare a “disegnare” le traiettorie. Ciò che si usa fare è eliminare la variabile t che, in fondo, ha un interesse marginale alla luce delle precedenti considerazioni: data una soluzione $\mathbf{y} = (y_1, y_2)$ della (1.59), si sceglie di disegnarne *l'immagine* nel piano (y_1, y_2) , che prende il nome di *piano delle fasi*³. In questo senso, la soluzione $\mathbf{y}$ può essere pensata come la traiettoria percorsa da un punto materiale che si muove nel piano (o, volendo essere più precisi, nel dominio Ω di $\mathbf{f}$). Inoltre, se l'origine dei tempi è $t_0 = 0$, si può evidenziare la “posizione” del punto nell'istante iniziale e indicare il verso di percorrenza. Il Teorema di esistenza e unicità garantisce che la traiettoria $\mathbf{y}$ non può “autointersecarsi” formando “occhielli”; altrimenti, troveremmo due soluzioni locali del problema di Cauchy ottenuto scegliendo come dato il punto di intersezione. Può invece succedere invece il fenomeno dell'*orbita periodica*:

Esercizio 1.36. Rappresentare nel piano (y_1, y_2) la soluzione del problema di Cauchy $y'_1 = -y_2$, $y'_2 = y_1$, $\mathbf{y}(0) = (1, 0)$.

Esercizio 1.37. Rappresentare nel piano delle fasi monodimensionale y (cioè nella retta reale) le soluzioni delle equazioni differenziali $y' = y$ e $y' = 1$ con la condizione iniziale (uguale per entrambe) $y(0) = 1$. Quale informazione si perde rispetto alla rappresentazione consueta nel piano (t, y) ?

Questo tipo di rappresentazione viene utilizzato anche per le equazioni scalari di ordine superiore; ad esempio, data un'equazione scalare autonoma del secondo ordine

³segnaliamo che anche per $N \geq 3$ si usa l'espressione *spazio delle fasi* per indicare lo spazio $(y_1, \dots, y_N)$, anche se naturalmente non possiamo “disegnarlo”

di incognita y , ponendo $v := y'$, si ottiene, come è noto, un sistema vettoriale del primo ordine nell'incognita (y, v) . In un'interpretazione cinematica, il piano delle fasi si ottiene ora mettendo in ascissa la *posizione* e in ordinata la *velocità* del punto che si muove in $\mathbb{R}$. Il tempo, anche qui, viene eliminato.

Infine, disegnare la traiettoria $\mathbf{y}$ nel piano delle fasi “partendo” dal tempo iniziale t_0 è particolarmente suggestivo quando si vuole studiare un problema di Cauchy in avanti per l'equazione (1.59) e si fa tendere il tempo a $+\infty$ (sempreché la soluzione massimale abbia come dominio tutto $\mathbb{R}^+$...). In effetti, sorge spontanea la domanda se $\mathbf{y}$ debba “arrestarsi” da qualche parte (ad esempio in un punto, o sul bordo di Λ), ovvero “esploda” all'infinito, o abbia un comportamento ancora più complicato. Tratteremo in dettaglio queste questioni nel terzo capitolo.

2 Sistemi lineari

2.1 Richiami di algebra lineare – diagonalizzazione

La trattazione dei sistemi di equazioni lineari del primo ordine richiede una certa dimestichezza con l'algebra lineare. In particolare, in questo paragrafo richiameremo brevemente alcune nozioni sulla diagonalizzazione delle matrici quadrate. La trattazione sarà completamente euristica e il lettore dovrebbe già sapere più di quanto scritto qui; tuttavia, questi richiami ci daranno modo di fissare le notazioni e le convenzioni che saranno adottate nel seguito. Manterremo lo stesso tipo di approccio anche nel paragrafo successivo, dove ci occuperemo della forma canonica di Jordan. Neanche in quel caso daremo dimostrazioni, ma cercheremo almeno di fornire un'idea della teoria sottostante presentando gli enunciati rigorosi dei teoremi e discutendone il significato. Per approfondimenti ulteriori, rinviamo ai testi [6] e [3, Appendice].

Nel seguito, $\mathcal{M}(\mathbb{R}^N)$ (rispettivamente $\mathcal{M}(\mathbb{C}^N)$) denoterà lo spazio vettoriale delle matrici reali (risp. complesse) quadrate di ordine N ; I , o I_N , sarà la matrice identità in $\mathcal{M}(\mathbb{R}^N)$. Cominciamo con una serie di definizioni e notazioni nel caso reale, osservando da subito che tutto si può ripetere, con ovvie modifiche, anche in quello complesso.

Per evitare complicazioni, cercheremo di vedere sempre i vettori di $\mathbb{R}^N$, che saranno indicati con lettere minuscole in grassetto, come N -uple di numeri (se uno vuole, le coordinate rispetto alla base canonica $\{\mathbf{e}^i\}$), ricorrendo il meno possibile all'uso di basi astratte di $\mathbb{R}^N$. Inoltre penseremo agli elementi di $\mathbb{R}^N$ come vettori colonna.

Data dunque $A \in \mathcal{M}(\mathbb{R}^N)$, l'endomorfismo di $\mathbb{R}^N$ associato ad A viene visto come l'applicazione che associa alla N -upla $\mathbf{x} = (x_j)$ la N -upla $\mathbf{y} = (y_i)$, ove $y_i = \sum_j a_{ij}x_j$. Più in dettaglio, se $A = (a_{ij})$, denoteremo con $\mathbf{a}_i$ (risp. $\mathbf{a}^j$) l' i -esimo vettore riga (risp. j -esimo vettore colonna) di A . Dati N vettori $\mathbf{a}^j \in \mathbb{R}^N$, $j = 1, \dots, N$, indicheremo con $(\mathbf{a}^1 \dots \mathbf{a}^N)$ la matrice avente i vettori $\mathbf{a}^j$ come *colonne*. Osserviamo che, se $M \in \mathcal{M}(\mathbb{R}^N)$, allora

$$M(\mathbf{a}^1 \dots \mathbf{a}^N) = (M\mathbf{a}^1 \dots M\mathbf{a}^N), \quad (2.1)$$

ove naturalmente $M\mathbf{a}^j$ indica il vettore (colonna) ottenuto come prodotto della matrice M per il vettore $\mathbf{a}^j$.

Siano ora $M_j \in \mathcal{M}(\mathbb{R}^{N_j})$, $j = 1, \dots, k$, k matrici quadrate di ordini N_j . Indicherò

con $M = \text{diag}(M_1, \dots, M_k)$ la matrice di ordine $N = \sum_{j=1}^k N_j$ data da

$$M = \begin{pmatrix} M_1 & & \\ & \ddots & \\ & & M_k \end{pmatrix} \quad (2.2)$$

ove le matrici M_j sono collocate sulla diagonale principale e si intende che gli elementi non scritti sono nulli (adotteremo tale convenzione anche nel seguito). Una matrice che ammette una decomposizione del tipo (2.2) con $k > 1$ verrà detta *diagonale a blocchi*. In particolare, se $\lambda_1, \dots, \lambda_N$ sono scalari, indicheremo con $m_{ij} = \text{diag}(\lambda_1, \dots, \lambda_N) \in \mathcal{M}(\mathbb{R}^N)$ la matrice diagonale tale che $m_{ij} = \lambda_i \delta_{ij}$, ove δ_{ij} sono i simboli di Kronecker. Per fare un esempio concreto, abbiamo

$$\text{diag} \left(\begin{pmatrix} 1 & 2 \\ 3 & 5 \end{pmatrix}, 2 \right) = \begin{pmatrix} 1 & 2 & 0 \\ 3 & 5 & 0 \\ 0 & 0 & 2 \end{pmatrix}.$$

Siano ora $M = \text{diag}(M_1, \dots, M_k)$, $S = \text{diag}(S_1, \dots, S_k)$ due matrici diagonali formate da blocchi dei medesimi ordini N_j , $j = 1, \dots, k$. Si verifica allora facilmente che

$$MS = \text{diag}(M_1 S_1, \dots, M_k S_k). \quad (2.3)$$

Le matrici diagonali a blocchi hanno una notevole importanza nelle applicazioni. Sia infatti M come in (2.2) e sia $\mathbf{v} \in \mathbb{R}^N$. Si verifica allora che

$$M\mathbf{v} = \begin{pmatrix} M_1 \mathbf{v}_1 \\ \vdots \\ M_k \mathbf{v}_k \end{pmatrix},$$

ove $\mathbf{v}_1, \dots, \mathbf{v}_k$ è la decomposizione di $\mathbf{v}$ in k vettori di lunghezze $N_1, \dots, N_k$, rispettivamente. Utilizzando un approccio più astratto, questo vuol dire che $\mathbb{R}^N$ si decompone come somma diretta di k sottospazi E_j di dimensioni N_j (il primo dei quali, E_1 , sarà ad esempio generato dai primi N_1 vettori della base canonica $\{\mathbf{e}^i\}$); inoltre, anche l'endomorfismo associato ad M si decompone in k endomorfismi di ciascuno degli spazi E_j . In tali condizioni, si dice che i sottospazi E_j sono M -invarianti. È evidente che una situazione come quella descritta è particolarmente fortunata; infatti, gli endomorfismi indotti da M sugli E_j si rappresentano *nella base canonica* proprio con la matrice M_j . Dunque, lo studio delle proprietà di M si riconduce immediatamente allo studio delle matrici M_j , che sono di ordine inferiore e dunque più facili da trattare. Volendo generalizzare, data una qualunque matrice quadrata A , può essere utile cercare una base di $\mathbb{R}^N$ in cui l'endomorfismo associato ad A assume la forma (2.2). Questo si può riformulare dicendo che si cerca una decomposizione $\mathbb{R}^N = E_1 \oplus \dots \oplus E_k$ in un numero $k > 1$ di sottospazi A -invarianti, ossia tali che $A(E_i) \subset E_i$.

Il caso migliore di tutti, corrispondente a $k = N$ e su cui ora diamo qualche richiamo, è quello delle matrici diagonalizzabili. Peraltro, a questo proposito, è più comodo dapprima ambientare la teoria generale in campo complesso, per poi tornare a discutere il caso reale. Sia dunque $M \in \mathcal{M}(\mathbb{C}^N)$. Ricordiamo che $\lambda \in \mathbb{C}$ si dice

autovalore di M se e solo se esiste $\mathbf{z} \in \mathbb{C} \setminus \{0\}$ tale che $M\mathbf{z} = \lambda\mathbf{z}$. Un tale vettore $\mathbf{z}$ si chiama allora *autovettore* della matrice M associato all'autovalore λ . L'insieme degli autovettori associati a λ risulta essere uno spazio vettoriale, detto *autospazio* associato a λ . Si ricorda che $\lambda \in \mathbb{C}$ è autovalore per M se e solo se è soluzione dell'equazione

$$\det(A - \lambda I) = 0 \quad (2.4)$$

ove $P_A(\lambda) := \det(A - \lambda I)$ viene ad essere un polinomio di grado N in λ , detto *polinomio caratteristico* della matrice A . Più in dettaglio, si dice *molteplicità algebrica* di un autovalore λ la sua molteplicità come zero del polinomio caratteristico.

Due proprietà importanti relative agli autovalori di una matrice sono le seguenti:

- due autovettori (non nulli) associati ad autovalori distinti di A risultano essere *linearmente indipendenti*;
- gli autospazi relativi all'endomorfismo A sono A -invarianti.

Per proseguire con la teoria, a questo punto è necessario richiamare qualcosa sui cambiamenti di coordinate. Anche qui, utilizzeremo un'impostazione “concreta”, evitando di ricorrere all'uso di basi astratte di $\mathbb{C}^N$. Dunque, supponiamo che $M \in \mathcal{M}(\mathbb{C}^N)$ sia una matrice **invertibile**. Chiameremo *cambiamento di coordinate* associato ad M l'endomorfismo di $\mathbb{C}^N$ che associa al generico vettore $\mathbf{u}$ il vettore $\mathbf{v} := M^{-1}\mathbf{u}$. Data una nuova, generica matrice $A \in \mathcal{M}(\mathbb{C}^N)$, risulta allora univocamente individuata $J \in \mathcal{M}(\mathbb{C}^N)$ tale che $A = MJM^{-1}$. Osserviamo che, $\forall \mathbf{u} \in \mathbb{C}^N$, si ha $A\mathbf{u} = MJ\mathbf{v}$, ossia

$$M^{-1}A\mathbf{u} = J\mathbf{v}. \quad (2.5)$$

Questa formula, che volendo ci si può limitare a interpretare come una regola di calcolo, da un punto di vista più astratto dice che l'endomorfismo associato ad A “opera” sulle nuove coordinate tramite la matrice J . In effetti, se $\mathbf{u}_1 := A\mathbf{u}$, $\mathbf{v}_1 := J\mathbf{v}$, allora $\mathbf{v}_1 = M^{-1}\mathbf{u}_1$; cioè il cambiamento di coordinate M trasforma $\mathbf{u}_1$ in $\mathbf{v}_1$. Dunque, il nostro obiettivo è quello di ricondurci al caso in cui J è più “semplice” possibile; in particolare, vogliamo vedere se e quando si può trovare un cambiamento di coordinate M che individui una J in forma diagonale.

Cominciamo con l'osservare che, se A, M, J verificano $A = MJM^{-1}$, allora i polinomi caratteristici di A e di J devono coincidere e dunque devono coincidere le loro radici, cioè gli autovalori. Dal momento che è chiaro che, se J è diagonale, i suoi autovalori sono gli elementi della diagonale, contati con le loro molteplicità, l'unica forma diagonale che possiamo sperare di ottenere per A è $\text{diag}(\lambda_1, \dots, \lambda_N)$, ove i λ_j sono i suoi autovalori, contati ciascuno con la sua molteplicità algebrica (in realtà l'unicità vale a meno di permutazioni dei λ_j). Come è noto, però, può capitare che A non sia *diagonalizzabile*, cioè che non esista alcun cambiamento di coordinate che fornisca la forma diagonale J .

Per capire il perchè partiamo comunque dal caso in cui A è diagonalizzabile e supponiamo di avere scritto $A = MJM^{-1}$ con J diagonale. Per quanto detto sopra, sappiamo allora come trovare J : basta calcolare gli autovalori di A . Per essere completamente soddisfatti, vogliamo però essere in grado di determinare anche la matrice M di cambiamento di coordinate. Scrivendo $J = \text{diag}(\lambda_1, \dots, \lambda_N)$, ove λ_i , $i = 1, \dots, N$

è un ordinamento qualunque degli autovalori di A contati con le loro molteplicità, consideriamo un vettore $\mathbf{e}^i$ della base canonica di $\mathbb{R}^N$ (visto come vettore colonna). Indicando con $\mathbf{m}^i$ i vettori colonna di M , otteniamo subito

$$A\mathbf{m}^i = A M \mathbf{e}^i = M J \mathbf{e}^i = M \lambda_i \mathbf{e}^i = \lambda_i M \mathbf{e}^i = \lambda_i \mathbf{m}^i.$$

Dunque, necessariamente, $\mathbf{m}^i$ è un autovettore associato all'autovalore λ^i . Viceversa, è altrettanto facile verificare che, se noi siamo in grado di trovare una matrice **invertibile** le cui colonne $\mathbf{m}^i$ sono autovettori di A associati agli autovalori λ_i , allora si ha che $A = M J M^{-1}$, ove $J = \text{diag}(\lambda_1, \dots, \lambda_N)$.

Da questo fatto, traiamo subito una conseguenza molto importante: se gli autovalori di A sono distinti, cioè hanno tutti molteplicità algebrica 1, allora A è senz'altro diagonalizzabile. Infatti, la definizione stessa implica che ad ogni autovalore è associato almeno un autovettore non nullo. Supponiamo invece che un autovalore λ abbia molteplicità algebrica $p > 1$. In questo caso, il discorso è più complicato. Chiamiamo *molteplicità geometrica* di λ la dimensione q dell'autospazio ad esso associato. La nostra speranza è che sia $p = q$, ma questo non è vero in generale. In effetti, si ha sempre che $1 \leq q \leq p$ (anche se la dimostrazione della seconda disuguaglianza non è immediata); può succedere però che $q < p$. In tal caso, non siamo in grado di esibire p autovettori linearmente indipendenti associati a λ e, di conseguenza, A non può essere diagonalizzabile. Riassumendo, A è diagonalizzabile se e solo se *la molteplicità algebrica di ogni autovalore è pari alla sua molteplicità geometrica*. Naturalmente, dovrebbe essere noto che esistono condizioni che assicurano la diagonalizzabilità di una matrice A (ad esempio se A è reale simmetrica). Per completezza, riportiamo alcuni esempi.

Esempio 2.1. Vediamo in concreto come si può stabilire se una matrice è diagonalizzabile e come si fa a determinare J e M . Consideriamo la matrice $A \in \mathcal{M}(\mathbb{R}^3)$ data da

$$A = \begin{pmatrix} 1 & 2 & 2 \\ 0 & 2 & 1 \\ 0 & -3 & -2 \end{pmatrix}.$$

Con un conto elementare si verifica che il polinomio caratteristico di A è

$$P_A(\lambda) = -(1 + \lambda)(1 - \lambda)^2.$$

Cominciamo a determinare l'autospazio relativo all'autovalore doppio $\lambda = 1$. Consideriamo un generico vettore (colonna) $\mathbf{v} = (a, b, c)$ e calcoliamo

$$(A - I)\mathbf{v} = \begin{pmatrix} 2b + 2c \\ b + c \\ -3b - 3c \end{pmatrix},$$

da cui l'unica condizione cui deve soddisfare $\mathbf{v}$ per essere autovettore è $b = -c$ e si possono facilmente esibire due autovettori linearmente indipendenti, ad esempio $\mathbf{v}^1 = (1, 1, -1)$ e $\mathbf{v}^2 = (0, 1, -1)$.

Imponendo invece $(A + I)\mathbf{v} = 0$, otteniamo il sistema

$$\begin{cases} 2a + 2b + 2c = 0 \\ 3b + c = 0 \\ -3b - c = 0, \end{cases}$$

Ora, sappiamo dalla teoria che lo spazio delle soluzioni dovrà avere dimensione 1. A noi, peraltro, basta calcolare una soluzione. Scegliendo, ad esempio, $a = 2$, otteniamo che un autovettore è $\mathbf{v}^3 = (2, 1, -3)$. Dunque, abbiamo che

$$AM = MJ, \quad \text{ove} \quad J = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{pmatrix}, \quad M = (\mathbf{v}^1 \ \mathbf{v}^2 \ \mathbf{v}^3) = \begin{pmatrix} 1 & 0 & 2 \\ 1 & 1 & 1 \\ -1 & -1 & -3 \end{pmatrix}$$

(chi non si fida faccia i conti).

Esempio 2.2. Sia ora invece

$$A = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 2 & 2 \end{pmatrix}$$

Procedendo come sopra, si vede che $P(\lambda) = (1 - \lambda)^2(2 - \lambda)$. Prendendo un generico vettore $\mathbf{v} = (a, b, c)$ e calcolando $(A - I)\mathbf{v}$, otteniamo il sistema

$$\begin{cases} b = 0 \\ 2b + c = 0, \end{cases}$$

da cui gli autovettori non nulli di A associati all'autovalore doppio $\lambda = 1$ sono tutti della forma $(a, 0, 0)$. Ne consegue che la molteplicità geometrica di $\lambda = 1$ è 1 e A dunque non è diagonalizzabile.

Esercizio 2.3. Sia $A = (a_{ij}) \in \mathcal{M}(\mathbb{C}^N)$. Supponiamo che M abbia un unico autovalore $\lambda \in \mathbb{C}$. Dimostrare che A è diagonalizzabile se e solo se è già nella forma scalare $A = \lambda I$.

Esercizio 2.4. Sia $A = (a_{ij}) \in \mathcal{M}(\mathbb{R}^3)$, definita da $a_{ij} = 1$ se $j = i + 1$ e $a_{ij} = 0$ altrimenti. Calcolare A^k per $k = 1, 2, 3$.

La teoria che abbiamo sviluppato fino a questo punto riguarda il caso complesso, anche se gli esempi che abbiamo dato sono ambientati in $\mathbb{R}^N$. Adattare tutto al caso reale non è peraltro una mera questione di notazioni. Infatti, se $A \in \mathcal{M}(\mathbb{R}^N)$, può ben succedere che il polinomio caratteristico della matrice A abbia radici complesse. Più in dettaglio, sappiamo dall'algebra che, se λ è un autovalore, allora anche il coniugato $\bar{\lambda}$ lo è. In effetti, se $\mathbf{w}$ è un autovettore associato a λ , si ha che

$$A\bar{\mathbf{w}} = \overline{A\mathbf{w}} = \overline{\lambda\mathbf{w}} = \bar{\lambda}\bar{\mathbf{w}};$$

dunque, $\bar{\mathbf{w}}$ è un autovettore associato a $\bar{\lambda}$ (il coniugato di un vettore di $\mathbb{C}^N$ si ottiene coniugando le singole componenti). Dunque, data $A \in \mathcal{M}(\mathbb{R}^N)$ diagonalizzabile, può darsi che sia la forma diagonale J , sia la matrice diagonalizzante M debbano essere necessariamente prese in $\mathcal{M}(\mathbb{C}^N)$. In effetti, in questa situazione si preferisce cercare una J un pochino diversa, come vedremo in un contesto più generale nel prossimo paragrafo.

2.2 Richiami di algebra lineare – forme canoniche

Data una matrice $A \in \mathcal{M}(\mathbb{R}^N)$, il nostro obiettivo iniziale era quello di trovare un cambiamento di coordinate M tale che $J = M^{-1}AM$ avesse una forma “più semplice” di A . Finora, però, abbiamo risolto il problema in modo in modo soddisfacente solo nel caso in cui A è diagonalizzabile e ha solo autovalori reali. Dunque, lo scopo di questo paragrafo è cercare dei cambiamenti di coordinate che consentano di trasformare **ogni** endomorfismo $A \in \mathcal{M}(\mathbb{R}^N)$ in una forma più semplice J , senza passare per il campo complesso. Naturalmente, le J che troveremo non potranno essere sempre diagonali; saranno, in generale, un pochino più complicate, ma comunque sufficientemente maneggevoli per le applicazioni che ci interessano. La base della teoria è costituita dal seguente teorema, piuttosto profondo, che non dimostriamo:

Teorema 2.5. *Sia $A \in \mathcal{M}(\mathbb{R}^N)$ e siano $\lambda_1, \dots, \lambda_k$ gli autovalori (distinti ed eventualmente complessi) di A . Allora, $\forall j = 1, \dots, k, \exists n_j \geq 1$ tale che*

$$E_j := \text{Ker}(A - \lambda_j I)^{n_j} = \text{Ker}(A - \lambda_j I)^{n_j+1}. \quad (2.6)$$

Inoltre, ponendo $A_j := A|_{E_j}$, si ha che esistono due applicazioni lineari Z_j, D_j di E_j in sè tali che

$$A_j = D_j + Z_j, \quad D_j Z_j = Z_j D_j, \quad D_j \mathbf{v} = \lambda_j \mathbf{v} \quad \forall \mathbf{v} \in E_j, \quad Z_j^{n_j} = 0. \quad (2.7)$$

Infine,

$$\mathbb{R}^N = E_1 \oplus \dots \oplus E_k. \quad (2.8)$$

La decomposizione (2.7) viene detta decomposizione primaria dell'operatore A . Lo spazio E_j si chiama autospazio generalizzato associato all'autovalore λ_j . D_j viene detta parte diagonale e Z_j parte nilpotente di A_j . ■

Vediamo in dettaglio che cosa significa il Teorema visto sopra, nel “caso modello” più facile di tutti: quello in cui A ha un solo autovalore reale λ . Abbiamo visto che A è diagonalizzabile se e solo se $\dim \text{Ker}(A - \lambda I) = N$, e in tal caso, il teorema precedente non dice nulla di nuovo. In caso contrario, $\text{Ker}(A - \lambda I)$ non contiene tutti i vettori di $\mathbb{R}^N$; tuttavia, se noi applichiamo ripetutamente $A - \lambda I$, prima o poi otteniamo l'operatore nullo. Abbiamo cioè la *catena di inclusioni*

$$\text{Ker}(A - \lambda I) \subset \text{Ker}(A - \lambda I)^2 \subset \dots \subset \text{Ker}(A - \lambda I)^n = \mathbb{R}^N$$

per qualche $n \geq 1$. Dunque, $\mathbb{R}^N$ non è l'autospazio associato a λ , ma, almeno, è l'autospazio generalizzato.

In generale, la situazione è più complessa e, in questo caso, la parte del teorema più importante per noi è quella che dice che, per ogni endomorfismo A di $\mathbb{R}^N$, è possibile trovare una base di $\mathbb{R}^N$ composta di autovettori generalizzati di A . Ora vediamo come questo risultato apparentemente molto astratto può essere usato per trovare una forma conveniente dell'endomorfismo A . In particolare, dapprima determineremo una matrice J simile ad A (ossia tale che $A = MJM^{-1}$), detta forma canonica di Jordan, e in un secondo momento vedremo anche come si fa a trovare la matrice M di cambiamento di coordinate.

Il punto di partenza è la (2.8). Questa formula assicura che, dato un qualsiasi $\mathbf{v} \in \mathbb{R}^N$, esistono (e sono unici) vettori $\mathbf{v}_j \in E_j$, $j = 1, \dots, k$, tali che $\mathbf{v} = \sum_j \mathbf{v}_j$. Inoltre,

$$A\mathbf{v} = A \sum_j \mathbf{v}_j = \sum_j A\mathbf{v}_j = \sum_j A_j \mathbf{v}_j.$$

In altre parole, l'endomorfismo A si decompone in k endomorfismi A_j che agiscono sugli autospazi generalizzati E_j (ciascuno dei quali è dunque A -invariante).

A questo punto, $\forall j$, sia $\mathbf{z}_{j,1}, \dots, \mathbf{z}_{j,N_j}$ una base di E_j (per il momento, non ci poniamo il problema della sua determinazione esplicita). Allora, i vettori

$$\mathbf{z}_{1,1}, \dots, \mathbf{z}_{1,N_1}, \dots, \mathbf{z}_{k,1}, \dots, \mathbf{z}_{k,N_k}$$

costituiscono una base di $\mathbb{R}^N$ nella quale l'endomorfismo A assume una forma diagonale a blocchi. Ammettiamo di esserci ricondotti a questo caso e, dunque, supponiamo che A sia già decomposto in questo modo: se riusciamo a trovare, per ogni blocco j , una forma semplice J_j per la matrice A , ossia un cambiamento di coordinate M_j , tale che $A_j = M_j J_j M_j^{-1}$, allora, grazie alla (2.3), si ha subito

$$A = M J M^{-1}, \quad \text{ove } J = \text{diag}(J_1, \dots, J_k), \quad M = \text{diag}(M_1, \dots, M_k)$$

e quindi ci siamo ricondotti al caso in cui A ha un solo autovalore λ , cosa che supponiamo da ora in avanti. Ora, in questo caso particolare, (2.7) ci dice che

$$A = D + Z, \quad \text{ove } D = \lambda I, \quad Z = A - \lambda I.$$

Supponendo che $M \in \mathcal{M}(\mathbb{R}^N)$ sia una qualunque matrice invertibile, si ha allora

$$M^{-1} A M = M^{-1} (D + Z) M = D + M^{-1} Z M;$$

infatti D è diagonale. Ossia, se noi operiamo un cambiamento di coordinate, questo muove solo la parte nilpotente di A ; è dunque sufficiente cercare una forma semplice per quest'ultima. La risposta è fornita dal seguente teorema, la cui dimostrazione, anch'essa non banale, viene omessa. Nell'enunciato, $\mathbf{0}^\sigma$ indica il vettore nullo di $\mathbb{R}^\sigma$.

Teorema 2.6. *Sia $Z_j \in \mathcal{M}(\mathbb{R}^N)$ una matrice nilpotente. Esistono allora matrici $G_1, \dots, G_r$, ove*

$$G_s \in \mathcal{M}^{t(s)+1}, \quad G_s = \begin{pmatrix} \mathbf{0}^{t(s)} & I_{t(s)} \\ 0 & \mathbf{0}_{t(s)} \end{pmatrix}, \quad t(s) \geq 0, \quad \forall s = 1, \dots, r. \quad (2.9)$$

ed esiste un cambiamento di coordinate M_j , tale che

$$Z_j = M_j G_j M_j^{-1}, \quad G_j = \text{diag}(G_1, \dots, G_r). \quad \blacksquare \quad (2.10)$$

La G_j è detta forma canonica nilpotente della matrice Z_j ed è unica a meno di permutazioni dei blocchi.

Per districarsi nella selva di indici dell'enunciato precedente è utile svolgere i seguenti tre esercizi, di carattere puramente tecnico:

Esercizio 2.7. Scrivere la forma esplicita di una matrice G della forma (2.9) e di ordine 1, 2, 3, rispettivamente.

Esercizio 2.8. Consideriamo la generica matrice nilpotente $Z \in \mathcal{M}(\mathbb{R}^4)$, decomposta come $Z = \text{diag}(G_1, \dots, G_r)$ con G_s della forma (2.9) $\forall s = 1, \dots, r$. Definiamo l'*indice di nilpotenza* di Z come il più piccolo numero naturale k tale che $Z^k = 0$. Quanto può valere al massimo k e come è fatta Z in questo caso? Come è fatta Z se $k = 1$? Esistono due matrici $Z_1, Z_2 \in \mathcal{M}(\mathbb{R}^4)$ *non simili* con indice di nilpotenza pari a 2 per entrambe? (Suggerimento: svolgere prima l'Esercizio 2.4)

Esercizio 2.9. Scrivere una forma canonica nilpotente della matrice

$$Z = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Proviamo a ricapitolare il discorso euristico fatto fino a questo momento: data una qualsiasi matrice A , abbiamo dapprima isolato i suoi autospazi generalizzati E_j ; operando sul singolo spazio E_j , abbiamo visto che $A|_{E_j}$, attraverso un nuovo cambiamento di coordinate può essere scritta come $D_j + G_j$, ove $D_j = \lambda_j I_{N_j}$ e G_j è della forma (2.10). Mettendo assieme il tutto,

Teorema 2.10. Sia $A \in \mathcal{M}(\mathbb{R}^N)$ e supponiamo che tutti i suoi autovalori siano reali. Allora esiste un cambiamento di coordinate M tale che $A = MJM^{-1}$, ove la forma canonica di Jordan J è una matrice diagonale a blocchi data da

$$J = \text{diag}(J_1, \dots, J_k), \quad (2.11)$$

ove ogni matrice elementare J_j è detta blocco di Jordan associato all'autovalore λ_j , ha dimensione N_j pari alla molteplicità algebrica di λ_j ed è della forma $J_j = D_j + G_j$, ove $D_j = \lambda_j I_{N_j}$ e G_j è come in (2.10). ■

Si noti che la molteplicità algebrica dell'autovalore λ_j è pari alla sua molteplicità geometrica se e solo se $G_j = 0$, ossia se la parte nilpotente di J_j è nulla.

Osservazione 2.11. Si noti che, “mettendo assieme” i contributi dei vari blocchi di Jordan, anche la matrice J si decompone come $J = D + G$, ove D è diagonale e G nilpotente. Si noti inoltre che $DG = GD$ (infatti è chiaro che commutano i singoli blocchi G_j e D_j ; inoltre possiamo usare la (2.3)).

Per concludere il discorso in modo utile per le applicazioni ci resta da dare una risposta alle seguenti questioni:

1. Il discorso fatto si ripete senza variazioni nel caso in cui A ha anche autovalori complessi. Tuttavia, in questo caso sia la forma canonica J che la matrice di trasformazione M saranno complesse. Come procedere?
2. Sappiamo che esiste una forma canonica J (che con poca fatica si può mostrare essere unica a meno di permutazioni di blocchi) e sappiamo come è fatta. Ma, invece, come è fatta la matrice di trasformazione M ? E come la si può calcolare?

Cominciamo dal primo problema. Dato che questo era rimasto in sospeso dal paragrafo precedente, iniziamo con l'analizzare un caso estremamente semplice, cioè quello di una $A \in \mathcal{M}(\mathbb{R}^2)$ diagonalizzabile con autovettori complessi coniugati $\boldsymbol{\mu} \pm i\boldsymbol{\nu}$ associati agli autovalori $a \pm ib$. Poiché una matrice (complessa) diagonalizzante è data da $M = (\boldsymbol{\mu} + i\boldsymbol{\nu} \quad \boldsymbol{\mu} - i\boldsymbol{\nu})$ ed essendo $AM = MD$, si ha

$$(A(\boldsymbol{\mu} + i\boldsymbol{\nu}) \quad A(\boldsymbol{\mu} - i\boldsymbol{\nu})) = ((a + ib)(\boldsymbol{\mu} + i\boldsymbol{\nu}) \quad (a - ib)(\boldsymbol{\mu} - i\boldsymbol{\nu})).$$

Andando quindi a calcolare le parti reale ed immaginaria, ad esempio, della prima colonna, otteniamo facilmente la formula

$$A(\boldsymbol{\nu} \quad \boldsymbol{\mu}) = (a\boldsymbol{\nu} + b\boldsymbol{\mu} \quad a\boldsymbol{\mu} - b\boldsymbol{\nu}) = (\boldsymbol{\nu} \quad \boldsymbol{\mu}) \begin{pmatrix} a & -b \\ b & a \end{pmatrix}, \quad (2.12)$$

da cui, ponendo

$$L := (\boldsymbol{\nu} \quad \boldsymbol{\mu}), \quad B := \begin{pmatrix} a & -b \\ b & a \end{pmatrix}, \quad (2.13)$$

vediamo che $A = LBL^{-1}$ con $L, B \in \mathcal{M}(\mathbb{R}^N)$. La matrice B , detta *forma canonica reale*, risulterà nelle applicazioni più comoda della forma diagonale complessa D .

Osservazione 2.12. Il lettore nota, o ricorda, qualche ulteriore analogia tra la matrice B e il numero complesso $a + ib$?

Diciamo due parole sull'estensione di questo trucco a situazioni più generali. Nel caso di una A diagonalizzabile che abbia tra gli autovalori i numeri complessi $a \pm ib$, con associati gli autovettori *linearmente indipendenti* $\boldsymbol{\mu}_1 \pm i\boldsymbol{\nu}_1, \dots, \boldsymbol{\mu}_k \pm i\boldsymbol{\nu}_k$, sostituendo nella matrice M di diagonalizzazione le colonne

$$\boldsymbol{\mu}_1 + i\boldsymbol{\nu}_1 \quad \boldsymbol{\mu}_1 - i\boldsymbol{\nu}_1 \quad \dots \quad \boldsymbol{\mu}_k + i\boldsymbol{\nu}_k \quad \boldsymbol{\mu}_k - i\boldsymbol{\nu}_k$$

con le colonne

$$\boldsymbol{\nu}_1 \quad \boldsymbol{\mu}_1 \quad \dots \quad \boldsymbol{\nu}_k \quad \boldsymbol{\mu}_k,$$

si otterrà una forma diagonale a blocchi, ove il blocco relativo agli autovalori coniugati $a \pm ib$ sarà della forma

$$\begin{pmatrix} B & & \\ & \ddots & \\ & & B \end{pmatrix},$$

ed ove ciascuno dei blocchi B di ordine 2 ha l'espressione data in (2.13). Infine, nel caso non diagonalizzabile, si può mostrare che è possibile ottenere una nuova forma H dall'espressione formalmente identica a quella data in (2.11), con la differenza che i blocchi H_j associati agli autovalori complessi coniugati si decomporranno in sottoblocchi $\tilde{G}_r$ della forma

$$\tilde{G}_r = \begin{pmatrix} B & I_2 & & \\ & B & \ddots & \\ & & \ddots & I_2 \\ & & & B \end{pmatrix},$$

ove I_2 è la matrice identica di ordine 2.

Veniamo infine alla determinazione della matrice di trasformazione M che porta A nella sua forma di *Jordan* (reale o complessa). Si può allora mostrare (e non è difficile, si procede esattamente come nel caso in cui A è diagonalizzabile, di cui questa è la logica estensione) che M ha come colonne gli *autovettori generalizzati* di A . Il problema pratico è come fare a determinare tali autovettori generalizzati, e questo può essere in generale abbastanza complesso, specie nel caso in cui N è grande. Invece, in dimensione bassa, di solito si riesce senza troppi problemi, ed il lettore è incoraggiato a svolgere qualche esercizio in proposito. Segnaliamo che un esempio svolto in dettaglio si trova in [3, pp. 692–693].

Esercizio 2.13. Se $A \in \mathcal{M}(\mathbb{R}^N)$, sarà vero che, se $\mathbf{m}$ è un autovettore generalizzato di A associato ad $a + ib$, allora $\overline{\mathbf{m}}$ è un autovettore generalizzato associato ad $a - ib$? Come sarà fatta la matrice N che trasforma A nella sua forma canonica “di Jordan reale” H ? (Vedi sotto per la risposta)

Riepilogo. Dal momento che le notazioni nel caso N -dimensionale sono piuttosto astratte e forse difficili da “digerire”, vediamo qualche esempio svolto in dettaglio per $N = 3, 4$, riferendoci al caso in cui è data una matrice A con un solo autovalore λ .

Sia, per cominciare, $A \in \mathcal{M}(\mathbb{R}^3)$ e $\lambda \in \mathbb{R}$. Allora la forma di Jordan di A può essere di tre tipi (a meno di permutazioni dei blocchi):

$$J_1 = \lambda I; \quad J_2 = \begin{pmatrix} \lambda & 1 & 0 \\ 0 & \lambda & 0 \\ 0 & 0 & \lambda \end{pmatrix}; \quad J_3 = \begin{pmatrix} \lambda & 1 & 0 \\ 0 & \lambda & 1 \\ 0 & 0 & \lambda \end{pmatrix}.$$

Nel caso J_1 , A è diagonalizzabile. Nel caso J_2 abbiamo due blocchi di Jordan di ordine 2 e 1, rispettivamente; la molteplicità geometrica di λ è 2 ($J_2 - \lambda I$ ha rango 1) e l'indice di nilpotenza di $J_2 - \lambda I$ (o, equivalentemente, di $A - \lambda I$) è 2. Nel caso J_3 c'è un unico blocco di Jordan di ordine 3; la molteplicità geometrica di λ è 1 ($J_3 - \lambda I$ ha rango 2) e l'indice di nilpotenza di $J_3 - \lambda I$ (o, equivalentemente, di $A - \lambda I$) è 3. In quest'ultimo caso, la matrice M di cambiamento di coordinate avrà come colonne un autovettore $\mathbf{m}^1$, un vettore $\mathbf{m}^2$ tale che $(A - \lambda)\mathbf{m}^2 = \mathbf{m}^1$ e un vettore $\mathbf{m}^3$ tale che $(A - \lambda)\mathbf{m}^3 = \mathbf{m}^2$.

Sia ora invece $A \in \mathcal{M}(\mathbb{R}^4)$ non diagonalizzabile e supponiamo che A abbia gli autovalori complessi coniugati $a \pm ib$ di molteplicità 2. Allora le forme di Jordan complessa J e reale H di A sono date da

$$J = \begin{pmatrix} a + ib & 1 & 0 & 0 \\ 0 & a + ib & 0 & 0 \\ 0 & 0 & a - ib & 1 \\ 0 & 0 & 0 & a - ib \end{pmatrix}, \quad H = \begin{pmatrix} a & -b & 1 & 0 \\ b & a & 0 & 1 \\ 0 & 0 & a & -b \\ 0 & 0 & b & a \end{pmatrix}.$$

Se $\mathbf{m}^1 = \boldsymbol{\mu}^1 + i\boldsymbol{\nu}^1$ è un autovettore associato a $a + ib$ e $\mathbf{m}^2 = \boldsymbol{\mu}^2 + i\boldsymbol{\nu}^2$ un autovettore generalizzato, allora le matrici M e N che portano A in J e, rispettivamente, in H , sono date da

$$M = (\mathbf{m}^1 \ \mathbf{m}^2 \ \overline{\mathbf{m}^1} \ \overline{\mathbf{m}^2}), \quad N = (\boldsymbol{\nu}^1 \ \boldsymbol{\mu}^1 \ \boldsymbol{\nu}^2 \ \boldsymbol{\mu}^2).$$

2.3 Teoria generale

In questo paragrafo daremo un'introduzione alla teoria generale dei sistemi di equazioni differenziali lineari del primo ordine, soffermandoci, più che sulle dimostrazioni, sulle notazioni e la terminologia che saranno usate diffusamente nel seguito.

Sia $I = (t_1, t_2) \subset \mathbb{R}$ un intervallo aperto, $A \in C^0(I; \mathcal{M}(\mathbb{R}^N))$, $\mathbf{b} \in C^0(I; \mathbb{R}^N)$. Queste ipotesi saranno conservate per tutto il paragrafo. Dati $t_0 \in I$ e $\mathbf{y}_0 \in \mathbb{R}^N$, consideriamo la versione lineare del problema di Cauchy in avanti (1.17), che assume la forma

$$\begin{cases} \mathbf{y}'(t) = A(t)\mathbf{y}(t) + \mathbf{b}(t) & \text{per } t \geq t_0 \\ \mathbf{y}(t_0) = \mathbf{y}_0. \end{cases} \quad (2.14)$$

Il precedente sistema si dice *omogeneo* se $\mathbf{b} \equiv 0$, cioè si ha

$$\begin{cases} \mathbf{y}'(t) = A(t)\mathbf{y}(t) & \text{per } t \geq t_0 \\ \mathbf{y}(t_0) = \mathbf{y}_0. \end{cases} \quad (2.15)$$

Vogliamo dimostrare che il teorema di esistenza e unicità in grande può essere applicato a (2.14), cosicché la soluzione massimale avrà tempo finale t_2 . Tuttavia, la verifica di questo fatto, abbastanza evidente nel caso monodimensionale, necessita nella situazione generale di un po' di lavoro poiché non è immediato mostrare la condizione di crescita (1.29). A questo scopo, richiamiamo qualche nozione sulle norme di matrici:

Definizione 2.14. Sia $\|\cdot\|$ una norma sullo spazio $\mathcal{M}(\mathbb{R}^N)$. Dico che $\|\cdot\|$ è una norma di matrici se vale la seguente proprietà:

$$\|AB\| \leq \|A\| \|B\| \quad \forall A, B \in \mathcal{M}(\mathbb{R}^N). \quad (2.16)$$

Lemma 2.15. La norma euclidea

$$\|\cdot\|_2 : \mathcal{M}(\mathbb{R}^N) \rightarrow \mathbb{R}, \quad \|A\|_2^2 := \sum_{i,j=1}^N a_{ij}^2 \quad (2.17)$$

è una norma di matrici su $\mathcal{M}(\mathbb{R}^N)$.

Prova. Siano $A, B \in \mathcal{M}(\mathbb{R}^N)$, $C = AB$. Siano a_{ij} , b_{ij} , c_{ij} , le componenti di A, B, C , rispettivamente. Allora,

$$c_{ij} = \sum_k a_{ik} b_{kj}, \quad \text{ossia} \quad c_{ij} = \mathbf{a}_i \cdot \mathbf{b}^j$$

e, per la disuguaglianza di Schwarz,

$$c_{ij}^2 \leq |\mathbf{a}_i|^2 |\mathbf{b}^j|^2 = \sum_k a_{ik}^2 \sum_k b_{kj}^2.$$

La tesi segue sommando rispetto a i e j . ■

Procedendo come nella dimostrazione precedente, si può anche provare che

$$|A\mathbf{v}| \leq \|A\|_2 |\mathbf{v}| \quad \forall A \in \mathcal{M}(\mathbb{R}^N), \quad \mathbf{v} \in \mathbb{R}^N. \quad (2.18)$$

Molto spesso viene usata una norma differente su $\mathcal{M}(\mathbb{R}^N)$, la cui definizione è legata alla proprietà (2.18) e all'interpretazione delle matrici come operatori lineari su $\mathbb{R}^N$.

Lemma 2.16. *L'applicazione*

$$\|\cdot\|_{\mathcal{L}} : \mathcal{M}(\mathbb{R}^N) \rightarrow \mathbb{R}, \quad \|A\|_{\mathcal{L}} := \sup \{|A\mathbf{v}|, \mathbf{v} \in \mathbb{R}^N, |\mathbf{v}| \leq 1\} \quad (2.19)$$

è una norma di matrici su $\mathcal{M}(\mathbb{R}^N)$.

Prova. Cominciamo col mostrare che $\|\cdot\|_{\mathcal{L}}$ è una norma. Sia $A \in \mathcal{M}(\mathbb{R}^N)$. È ovvio dalla definizione che $\|A\|_{\mathcal{L}} \geq 0$. Se A non è la matrice nulla, certamente esiste un vettore $\mathbf{v} \in \mathbb{S}^{N-1}$ tale che $A\mathbf{v} \neq 0$. Dunque $\|A\|_{\mathcal{L}} > 0$. Dimostrare che $\|\lambda A\|_{\mathcal{L}} = |\lambda| \|A\|_{\mathcal{L}}$ è semplice. Infine, la disuguaglianza triangolare segue dal fatto che, per ogni $\mathbf{v} \in \mathbb{R}^N$ con $|\mathbf{v}| \leq 1$, si ha

$$|(A+B)\mathbf{v}| = |A\mathbf{v} + B\mathbf{v}| \leq |A\mathbf{v}| + |B\mathbf{v}| \leq \|A\|_{\mathcal{L}} + \|B\|_{\mathcal{L}}$$

e basta prendere il sup a primo membro. Infine, la proprietà (2.16) è un'immediata conseguenza del fatto che

$$|A\mathbf{v}| \leq \|A\|_{\mathcal{L}} |\mathbf{v}| \quad \forall A \in \mathcal{M}(\mathbb{R}^N), \quad \mathbf{v} \in \mathbb{R}^N. \quad \blacksquare \quad (2.20)$$

Ricordiamo ora il

Teorema 2.17. *Sia X uno spazio vettoriale reale di dimensione finita e siano $\|\cdot\|_1$ e $\|\cdot\|_2$ norme su X . Allora, $\|\cdot\|_1$ e $\|\cdot\|_2$ sono equivalenti, ossia esistono $c_1, c_2 > 0$ tali che*

$$\|v\|_1 \leq c_1 \|v\|_2, \quad \|v\|_2 \leq c_2 \|v\|_1, \quad \forall v \in X. \quad \blacksquare \quad (2.21)$$

Grazie a questo risultato, ogni norma $\|\cdot\|$ sullo spazio $\mathcal{M}(\mathbb{R}^N)$ verifica, per qualche $c \geq 1$ (indipendente da A, B),

$$\|AB\| \leq c \|A\| \|B\| \quad \forall A, B \in \mathcal{M}(\mathbb{R}^N), \quad (2.22)$$

che può essere vista come una versione debole della proprietà (2.16). Notiamo infine che da (2.18) e (2.19) segue

$$\|A\|_{\mathcal{L}} \leq \|A\|_2 \quad \forall A \in \mathcal{M}(\mathbb{R}^N), \quad (2.23)$$

ove la disuguaglianza può essere stretta. Ad esempio, se I è la matrice identica, si ha $\|I\|_{\mathcal{L}} = 1$ e $\|I\|_2 = N^{1/2}$.

Corollario 2.18. *Il Problema (2.14) ammette una e una sola soluzione massimale “in grande” $(\mathbf{y}, t_2)$.*

Prova. Vogliamo naturalmente applicare il Teorema 1.23, di cui verifichiamo le ipotesi. Innanzitutto, abbiamo che $\mathbf{f}(t, \mathbf{y}) = A(t)\mathbf{y} + \mathbf{b}(t)$, da cui, sicuramente, $\mathbf{f} \in C^0(I \times \mathbb{R}^N; \mathbb{R}^N)$. Si ha inoltre che

$$|\mathbf{f}(t, \mathbf{y}_1) - \mathbf{f}(t, \mathbf{y}_2)| = |A(t)(\mathbf{y}_1 - \mathbf{y}_2)| \leq \|A(t)\|_{\mathcal{L}} |\mathbf{y}_1 - \mathbf{y}_2|,$$

cosicché (1.34) è verificata con la scelta di $L(t) = \|A(t)\|_{\mathcal{L}}$ (si ricordi che la norma è una funzione continua, cfr. (1.1)). $\blacksquare$

Dal momento che la soluzione massimale del problema (2.14) è definita su tutto I , d'ora in poi, quando parleremo della “soluzione” di un sistema lineare, intenderemo sempre la soluzione in grande, che sarà indicata senza specificarne il dominio. Peraltro, per affrontare il problema della risoluzione esplicita di (2.14), conviene inizialmente “dimenticarsi” della condizione iniziale e studiare la struttura dell'insieme delle soluzioni dell'equazione lineare

$$\mathbf{y}'(t) = A(t)\mathbf{y}(t) + \mathbf{b}(t). \quad (2.24)$$

Chiamiamo naturalmente *equazione omogenea* associata a (2.24) la

$$\mathbf{y}'(t) = A(t)\mathbf{y}(t). \quad (2.25)$$

Vale il seguente teorema fondamentale, la cui dimostrazione, peraltro già nota al lettore, si trova in [4, Teorema IX.4.1] (cui si rinvia per un enunciato più dettagliato):

Teorema 2.19. *L'insieme delle soluzioni dell'equazione (2.25) è uno spazio vettoriale $V \subset C^1(I; \mathbb{R}^N)$ di dimensione N . Inoltre, se $\bar{\mathbf{y}}$ è una qualsiasi soluzione dell'equazione (2.24), la totalità delle soluzioni di (2.24) è costituita dalla varietà affine $V + \bar{\mathbf{y}}$. Infine, $\forall t_0 \in I$, l'applicazione*

$$V \rightarrow \mathbb{R}^N, \quad \mathbf{y} \mapsto \mathbf{y}(t_0)$$

è un isomorfismo di spazi vettoriali. ■

Dunque, per l'integrazione dell'equazione lineare (2.24) si procede in generale in due passi: dapprima si determina lo spazio V , che comprende *tutte* le soluzioni di (2.25). Questo equivale naturalmente a trovare N soluzioni linearmente indipendenti, ossia una *base* di V . Quindi, si cerca *una* soluzione $\bar{\mathbf{y}}$ dell'equazione *completa* (2.24). Nel caso in cui si debba determinare *la* soluzione di un problema di Cauchy (o di altro tipo), il terzo passo del procedimento consisterà naturalmente nell'imporre le condizioni iniziali per determinare il valore di N costanti arbitrarie. In effetti, grazie all'ultima parte dell'enunciato, se $\mathbf{y}^1, \dots, \mathbf{y}^N$ sono soluzioni di (2.25), allora

$$\text{rango}(\mathbf{y}^1(t) \dots \mathbf{y}^N(t)) = \text{rango}(\mathbf{y}^1(t_0) \dots \mathbf{y}^N(t_0)) \quad \forall t \in I.$$

In particolare, se i vettori $\mathbf{y}^1(t_0), \dots, \mathbf{y}^N(t_0)$, corrispondenti a N scelte dei dati iniziali, sono linearmente indipendenti, allora risultano linearmente indipendenti i vettori $\mathbf{y}^1(t), \dots, \mathbf{y}^N(t)$ per ogni $t \in I$. In questo caso, la matrice $S(t) := (\mathbf{y}^1(t) \dots \mathbf{y}^N(t))$ prende il nome di *matrice fondamentale* del sistema (2.25). La notazione matriciale è particolarmente comoda per vari motivi. Ad esempio, si vede subito che, se $\mathbf{c} \in \mathbb{R}^N$, allora la generica soluzione di (2.25) può essere scritta come $S(t)\mathbf{c}$, dove $\mathbf{c}$ si interpreta come un vettore di costanti arbitrarie. Inoltre, la matrice $S(t)$ verifica l'equazione differenziale matriciale

$$S'(t) = A(t)S(t). \quad (2.26)$$

Osserviamo infine che, naturalmente, la matrice S non è unica: dati N vettori linearmente indipendenti $\mathbf{u}^i \in \mathbb{R}^N$ e posto $U := (\mathbf{u}^1 \dots \mathbf{u}^N)$, si ha che $T := SU$ è ancora una matrice fondamentale. Le verifiche di questi fatti, basate essenzialmente sulla proprietà (2.1), sono lasciate al lettore.

Se si sta analizzando un Problema di Cauchy e dunque è assegnato un tempo iniziale $t_0 \in I$, per comodità di calcolo spesso conviene considerare una particolare matrice fondamentale, ossia quella (chiamiamola R) che verifica $R(t_0) = I$ (il lettore si chieda perché R è unica). Questa matrice si chiama *matrice risolvente* del sistema; si usa più spesso nel caso in cui A è a coefficienti costanti e $t_0 = 0$. Data una *qualsiasi* matrice fondamentale $S(t)$, la matrice risolvente si calcola facilmente essendo $R(t) = S(t)S^{-1}(0)$.

Variazione delle costanti. Vogliamo ora affrontare il problema della determinazione esplicita delle soluzioni di (2.24). Grazie al Teorema 2.19, dobbiamo fare due cose: trovare lo spazio V e trovare una soluzione $\bar{\mathbf{y}}$. Lasciamo per il seguito il primo problema, che è il più difficile, e ci occupiamo per il momento del secondo. Supponiamo dunque di essere così fortunati da conoscere una matrice fondamentale $S(t)$ di (2.25). Vale allora il

Teorema 2.20. *Esiste una soluzione $\bar{\mathbf{y}}(t)$ di (2.24) della forma $\bar{\mathbf{y}}(t) = S(t)\mathbf{w}(t)$ per un'opportuna scelta della funzione $\mathbf{w} \in C^1(I; \mathbb{R}^N)$.*

Prova. Sostituisco formalmente $\bar{\mathbf{y}}(t) = S(t)\mathbf{w}(t)$ in (2.24). Imponendo che $\bar{\mathbf{y}}$ sia soluzione e usando la (2.26), ho subito

$$A\bar{\mathbf{y}} + \mathbf{b} = \bar{\mathbf{y}}' = (S\mathbf{w})' = S'\mathbf{w} + S\mathbf{w}' = AS\mathbf{w} + S\mathbf{w}' = A\bar{\mathbf{y}} + S\mathbf{w}',$$

da cui ottengo $\mathbf{b} = S\mathbf{w}'$ e

$$\mathbf{w}' = S^{-1}\mathbf{b}. \quad (2.27)$$

A questo punto, il procedimento euristico viene reso rigoroso, ripercorrendo i passaggi alla rovescia e mostrando che ogni funzione $\mathbf{w}$ che verifica la (2.27) dà luogo a una soluzione $\bar{\mathbf{y}} = S\mathbf{w}$ di (2.24). In particolare, dato $t_0 \in I$, posso scegliere $\mathbf{w}$ in modo tale che sia $\mathbf{w}(t_0) = 0$. Questo fornisce subito la *formula risolutiva*

$$\bar{\mathbf{y}}(t) = S(t) \int_{t_0}^t S^{-1}(\tau)\mathbf{b}(\tau) d\tau. \quad \blacksquare \quad (2.28)$$

Dunque, la generica soluzione di (2.24) è data da

$$\mathbf{y}(t) = S(t)\mathbf{c} + S(t) \int_{t_0}^t S^{-1}(\tau)\mathbf{b}(\tau) d\tau, \quad \mathbf{c} \in \mathbb{R}^N. \quad (2.29)$$

Infine, se consideriamo il problema di Cauchy (2.14), sostituendo la condizione iniziale otteniamo che la soluzione è data da

$$\mathbf{y}(t) = S(t)S^{-1}(t_0)\mathbf{y}_0 + S(t) \int_{t_0}^t S^{-1}(\tau)\mathbf{b}(\tau) d\tau, \quad (2.30)$$

formula che ovviamente diventa più semplice qualora sia $S(t_0) = I$.

La precedente formula risolutiva fornisce l'espressione esplicita della soluzione. Tuttavia, naturalmente, nei casi pratici l'uso della (2.30) comporta calcoli molto pesanti,

senza contare che non sempre è possibile determinare $S(t)$. Una situazione in cui le cose risultano abbastanza semplici è il caso scalare:

$$y'(t) = a(t)y(t) + b(t), \quad y(t_0) = y_0, \quad (2.31)$$

ove sono date $a, b \in C^0(I; \mathbb{R})$. In tal caso, ponendo

$$A(t) := \int_{t_0}^t a(\tau) d\tau,$$

si vede facilmente che la (2.30) diventa

$$y(t) = y_0 e^{A(t)} + e^{A(t)} \int_{t_0}^t e^{-A(\tau)} b(\tau) d\tau. \quad (2.32)$$

Infatti, si ha che $S(t) = \exp(A(t))$ è la “matrice” risolvente. Ovviamente, è anche possibile scrivere una formula esplicita per $\bar{y}$ analoga a (2.28) e valida per il caso scalare.

Equazioni lineari di ordine superiore al primo. Come nel caso generale, la teoria delle equazioni lineari di ordine superiore al primo nella forma normale

$$y^{(k)}(t) = \sum_{i=0}^{k-1} a_i(t) y^{(i)}(t) + b(t), \quad (2.33)$$

ove $a_i, b \in C^0(I; \mathbb{R})$, si riconduce alla teoria dei sistemi del primo ordine tramite la sostituzione

$$\mathbf{y} := (y_1, \dots, y_k), \quad \text{ove } y_1 := y, \quad y_2 := y', \quad \dots, \quad y_k := y^{(k-1)}. \quad (2.34)$$

Infatti, sotto queste notazioni, (2.33) si riscrive in forma di sistema come

$$\mathbf{y}'(t) = A(t)\mathbf{y}(t) + \mathbf{b}(t), \quad \text{ove } A(t) = \begin{pmatrix} \mathbf{0}^{k-1} & I_{k-1} \\ a_0 & a_1 \dots a_k \end{pmatrix}, \quad \mathbf{b}(t) = \begin{pmatrix} \mathbf{0}^{k-1} \\ b(t) \end{pmatrix}. \quad (2.35)$$

Si ricorda che $\mathbf{0}^{k-1} \in \mathbb{R}^{k-1}$ è il vettore (colonna) nullo e I_{k-1} è la matrice identica di ordine $k-1$.

Dunque, è chiaro che, se noi siamo in grado di determinare l'insieme delle soluzioni di (2.35), allora, prendendo la *prima componente* di ognuna di queste, otterremo l'insieme delle soluzioni di (2.33). In particolare, il Teorema 2.19 ci assicura che le soluzioni di (2.33) costituiscono una varietà affine $V + \bar{y}$ di $C^1(I; \mathbb{R})$, ove $\bar{y}$ è una soluzione di (2.33) e V è lo spazio vettoriale, di *dimensione* k , formato dalle soluzioni dell'equazione omogenea associata a (2.33) (ossia dalle prime componenti delle soluzioni del sistema omogeneo associato a (2.35)). Notiamo inoltre che con un facile argomento di *bootstrap* si mostra che ogni soluzione $y \in V + \bar{y}$ ha la regolarità $C^k(I; \mathbb{R})$. Più in dettaglio, se $a_0, \dots, a_{k-1}, b \in C^{k_0}(I)$, per qualche $k_0 \geq 0$, allora y soddisfa $y \in C^{k_0+k}(I; \mathbb{R})$.

La teoria dei sistemi del primo ordine si trasporta quindi senza troppe modifiche al caso delle equazioni lineari di ordine $k \geq 1$. Val la pena comunque osservare qualcosa

sul metodo della variazione delle costanti. Innanzitutto, se $y^1, \dots, y^k$ sono k soluzioni linearmente indipendenti dell'equazione omogenea associata a (2.33), allora possiamo costruire una matrice fondamentale per il sistema associato:

$$S(t) = \begin{pmatrix} y^1(t) & \dots & y^k(t) \\ \vdots & & \vdots \\ (y^1)^{(k-1)}(t) & \dots & (y^k)^{(k-1)}(t) \end{pmatrix}. \quad (2.36)$$

A questo punto, considerando il sistema (2.35), possiamo cercarne una soluzione $\bar{\mathbf{y}}(t)$ della forma $\bar{\mathbf{y}}(t) = S(t)\mathbf{w}(t)$ ed arriviamo, come nel caso precedente, a $\mathbf{w}' = S^{-1}\mathbf{b}$. Tuttavia, nei casi pratici può non essere necessario calcolare la matrice inversa S^{-1} , ma procedere in modo più diretto. La situazione si chiarisce meglio tramite un

Esempio 2.21. Consideriamo il problema di Cauchy dato dall'equazione $y'' - 2y' + y = 2e^t$ con le condizioni $y(0) = 0$, $y'(0) = 1$. Si può dimostrare (vedi § 2.6) che $V = \text{span}\{e^t, te^t\}$. Allora, una matrice fondamentale è

$$S(t) = \begin{pmatrix} e^t & te^t \\ e^t & (t+1)e^t \end{pmatrix} = e^t \begin{pmatrix} 1 & t \\ 1 & t+1 \end{pmatrix}.$$

Applicando il metodo della variazione delle costanti, otteniamo

$$e^t \begin{pmatrix} 1 & t \\ 1 & t+1 \end{pmatrix} \begin{pmatrix} w'_1 \\ w'_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 2e^t \end{pmatrix},$$

da cui il sistema

$$w'_1 + tw'_2 = 0, \quad w'_1 + (t+1)w'_2 = 2$$

la cui soluzione nulla in 0 è $w_1 = -t^2$, $w_2 = 2t$. Per ottenere $\bar{\mathbf{y}}$, basta ora calcolare il *primo termine* del prodotto $\bar{\mathbf{y}} = S\mathbf{w}$.

Sia nel caso delle equazioni che in quello dei sistemi, resta però aperto il problema della determinazione dello spazio V , cioè dell'integrale generale dell'equazione omogenea. Purtroppo siamo in grado di dare una risposta completa a tale questione solo nel caso dei *coefficienti costanti*, ossia quando $\exists A \in \mathcal{M}(\mathbb{R}^N)$ tale che $A(t) \equiv A$ per ogni t (cfr. (2.25)), o una condizione analoga per (l'omogenea associata a) (2.33). Osserviamo che un sistema lineare omogeneo è a coefficienti costanti se e solo se è autonomo e, in tal caso, valgono le considerazioni di § 1.6. Nei prossimi tre paragrafi forniremo una risposta costruttiva a questa domanda, dapprima nel caso, più complicato, dei sistemi, quindi in quello della singola equazione di ordine k . Per il momento concludiamo il paragrafo con un trucco che può essere utile in certi casi particolari.

Abbassamento di ordine. Supponiamo di dover studiare l'equazione omogenea associata alla (2.33), cioè

$$y^{(k)}(t) = \sum_{i=0}^{k-1} a_i(t)y^{(i)}(t). \quad (2.37)$$

Se, per caso (ad esempio usando delle soluzioni di prova), siamo riusciti a determinare una soluzione particolare z della (2.37), allora possiamo ricondurci a un'equazione di

ordine $k-1$. Cerchiamo infatti una soluzione della forma $y = zv$. Grazie alla formula di Leibniz, abbiamo a primo membro

$$y^{(k)} = \sum_{r=0}^k \binom{k}{r} z^{(r)} v^{(k-r)} = z^{(k)} v + \text{altri termini.} \quad (2.38)$$

Analogamente, il secondo membro di (2.37) diventa

$$\sum_{i=0}^{k-1} a_i y^{(i)} = \sum_{i=0}^{k-1} a_i \sum_{s=0}^i \binom{i}{s} z^{(s)} v^{(i-s)} = \sum_{i=0}^{k-1} a_i z^{(i)} v + \text{altri termini.} \quad (2.39)$$

Eguagliando le due formule precedenti (imponendo cioè che y sia soluzione), si cancellano i termini ove compare v , cioè quelli scritti esplicitamente nei secondi membri. Gli “altri termini” rimasti contengono le derivate di z , che sono funzioni note, e le derivate di v dal primo al k -esimo ordine. Ne risulta dunque un’equazione lineare di ordine $k-1$, in genere non in forma normale, nell’incognita v' . In particolare, se $k=2$, possiamo usare la formula risolutiva (2.32) e determinare v' . Quindi si trovano v per integrazione e infine y ; una base di V sarà allora data da $\{z, y\}$.

2.4 Sistemi a coefficienti costanti diagonalizzabili

Supponiamo $A \in \mathcal{M}(\mathbb{R}^N)$. Vogliamo descrivere in dettaglio la struttura delle soluzioni del sistema lineare omogeneo a coefficienti costanti

$$\mathbf{y}' = A\mathbf{y}. \quad (2.40)$$

Grazie al Corollario 2.18, sappiamo che tale sistema ammette esattamente N soluzioni linearmente indipendenti $\mathbf{y}^1, \dots, \mathbf{y}^N$ definite su tutto $\mathbb{R}$. Il problema è il calcolo esplicito di tali soluzioni:

Esempio 2.22. Consideriamo il sistema *disaccoppiato*

$$y_1' = 4y_1, \quad y_2' = -3y_2, \quad (2.41)$$

Vogliamo calcolarne la generica soluzione $\mathbf{y} = (y_1, y_2)$. Osserviamo che tale sistema si riscrive in forma matriciale come

$$\mathbf{y}' = \begin{pmatrix} 4 & 0 \\ 0 & -3 \end{pmatrix} \mathbf{y}, \quad \text{ossia} \quad \begin{pmatrix} y_1' \\ y_2' \end{pmatrix} = \begin{pmatrix} 4 & 0 \\ 0 & -3 \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}.$$

Risolvendo le equazioni in (2.41), è facile vedere che la generica prima componente di una soluzione è della forma $y_1 = c_1 e^{4t}$ e la generica seconda componente è della forma $y_2 = c_2 e^{-3t}$, ove c_1, c_2 sono costanti arbitrarie. In forma vettoriale, possiamo scrivere che una base dello spazio delle soluzioni di (2.41) è data da

$$\mathbf{y}^1 = \begin{pmatrix} e^{4t} \\ 0 \end{pmatrix}, \quad \mathbf{y}^2 = \begin{pmatrix} 0 \\ e^{-3t} \end{pmatrix}.$$

Questo esempio ci permette di specificare un po' di notazioni: si osservi che le generiche soluzioni del sistema vettoriale sono state indicate coi *vettori colonna* $\mathbf{y}^j$, $j = 1, 2$. Invece la notazione y_i , con l'indice i in basso, indica le *componenti* di una certa soluzione $\mathbf{y}$. Ad esempio, le N soluzioni $\mathbf{y}^j$ si indicheranno esplicitamente come

$$\mathbf{y}^1 = \begin{pmatrix} y_1^1 \\ \vdots \\ y_N^1 \end{pmatrix} = \begin{pmatrix} y_{11} \\ \vdots \\ y_{1N} \end{pmatrix}, \quad \dots, \quad \mathbf{y}^N = \begin{pmatrix} y_1^N \\ \vdots \\ y_N^N \end{pmatrix} = \begin{pmatrix} y_{N1} \\ \vdots \\ y_{NN} \end{pmatrix}.$$

Tuttavia, se A non è in forma diagonale, descrivere lo spazio delle soluzioni di (2.40) non è semplice come nell'esempio precedente. L'osservazione fondamentale è contenuta nella seguente proposizione, di dimostrazione immediata:

Proposizione 2.23. *Sia $A \in \mathcal{M}(\mathbb{R}^N)$. Supponiamo che esistano due matrici $M, J \in \mathcal{M}(\mathbb{R}^N)$, con M invertibile, tali che $A = MJM^{-1}$. Allora, il cambiamento di variabile $\mathbf{y} = M\mathbf{z}$ trasforma il sistema (2.40) nel sistema equivalente*

$$\mathbf{z}' = J\mathbf{z}. \quad (2.42)$$

Dove sta il vantaggio? Ovviamente, in generale si spera che la nuova matrice J sia “più semplice” di A e, dunque, di poter descrivere con poca fatica una base di soluzioni $(\mathbf{z}^1, \dots, \mathbf{z}^N)$ per (2.42). Tra l'altro, per tornare alle “vecchie” coordinate $\mathbf{y}$, non è necessario calcolare la matrice inversa M^{-1} . Ponendo infatti $Z := (\mathbf{z}^1 \dots \mathbf{z}^N)$, è chiaro che Z è una matrice fondamentale per il sistema trasformato (2.42) e quindi

$$Y = (\mathbf{y}^1 \dots \mathbf{y}^N) := MZ = (M\mathbf{z}^1 \dots M\mathbf{z}^N),$$

è una matrice fondamentale per il sistema di partenza (2.40), dato che M è invertibile. Naturalmente, qualora si stia studiando un problema di Cauchy, la condizione da accoppiare all'equazione trasformata (2.42) sarà $\mathbf{z}(t_0) = \mathbf{z}_0 := M^{-1}\mathbf{y}_0$.

Chiaramente il caso più semplice, di cui ci occupiamo in questo paragrafo, è quello in cui A è diagonalizzabile. Vediamo un esempio:

Esempio 2.24. Consideriamo il sistema

$$\begin{cases} y_1' = 2y_1 + y_2 \\ y_2' = 12y_1 + 3y_2. \end{cases}$$

Diagonalizzando la matrice A dei coefficienti, si trova subito che

$$AM = MJ, \quad \text{ove} \quad J = \begin{pmatrix} 6 & 0 \\ 0 & -1 \end{pmatrix}, \quad M = \begin{pmatrix} 1 & 1 \\ 4 & -3 \end{pmatrix}.$$

Nelle nuove coordinate $\mathbf{z} = M^{-1}\mathbf{y}$, il sistema risulta disaccoppiato; infatti si decompone in

$$\begin{cases} z_1' = 6z_1, \\ z_2' = -z_2, \end{cases}$$

da cui $z_1 = c_1 e^{6t}$, $z_2 = c_2 e^{-t}$ e quindi una base di soluzioni del sistema è

$$\mathbf{z}^1 = \begin{pmatrix} e^{6t} \\ 0 \end{pmatrix}, \quad \mathbf{z}^2 = \begin{pmatrix} 0 \\ e^{-t} \end{pmatrix};$$

infine, una base di soluzioni nelle vecchie coordinate è data da

$$\mathbf{y}^1 = M \begin{pmatrix} e^{6t} \\ 0 \end{pmatrix} = e^{6t} \mathbf{m}^1, \quad \mathbf{y}^2 = M \begin{pmatrix} 0 \\ e^{-t} \end{pmatrix} = e^{-t} \mathbf{m}^2.$$

In generale, se $A = MJM^{-1}$ con $J = \text{diag}(\lambda_1, \dots, \lambda_N)$, il procedimento descritto fornisce subito la matrice risolvete del sistema trasformato (2.42) (relativa al tempo iniziale $t_0 = 0$), ossia

$$Z(t) = (e^{\lambda_1 t} \mathbf{e}^1 \dots e^{\lambda_N t} \mathbf{e}^N). \quad (2.43)$$

Dunque, una matrice fondamentale di (2.40) è

$$Y(t) = MZ(t) = (e^{\lambda_1 t} \mathbf{m}^1 \dots e^{\lambda_N t} \mathbf{m}^N). \quad (2.44)$$

Tale matrice non è però la matrice risolvete di (2.40), che ha l'espressione $S(t) = Y(t)M^{-1}$, come si verifica facilmente. Naturalmente, nella pratica, il calcolo esplicito di $S(t)$ richiede un certo numero di conti. La discussione fatta fino a questo punto copre il caso in cui A è diagonalizzabile e tutti i suoi autovalori sono reali. Nel caso vi siano autovalori complessi, il procedimento usato si può ripercorrere, ma il risultato non è del tutto soddisfacente, in quanto vengono necessariamente a comparire soluzioni complesse di sistemi a coefficienti reali, quando la teoria astratta garantisce che una base di soluzioni reali esiste sicuramente.

Esempio 2.25. Consideriamo il sistema (2.40) con la scelta

$$A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}, \quad \text{corrispondente a} \quad \begin{cases} y_1' = -y_2 \\ y_2' = y_1. \end{cases}$$

Allora, $P_A(\lambda) = \lambda^2 + 1$ ha le radici complesse $\pm i$. Osservo tuttavia che il sistema si può riscrivere facilmente come

$$\begin{cases} y_1'' = -y_1 \\ y_2 = -y_1'. \end{cases}$$

Dunque, la prima equazione, risolta esplicitamente, fornisce $y_1 = c_1 \cos t + c_2 \sin t$, da cui $y_2 = c_1 \sin t - c_2 \cos t$. In forma vettoriale,

$$\mathbf{y}^1 = \begin{pmatrix} \cos t \\ \sin t \end{pmatrix}, \quad \mathbf{y}^2 = \begin{pmatrix} \sin t \\ -\cos t \end{pmatrix}.$$

La situazione riportata in questo esempio è in un certo senso standard, in quanto è possibile appellarsi alla teoria svolta nel Paragrafo 2.2. In effetti, supponendo B come in (2.13), è facile vedere che una base di soluzioni del sistema 2×2

$$\mathbf{y}' = B\mathbf{y} \quad \text{è data da} \quad \mathbf{y}^1 = \begin{pmatrix} e^{at} \cos bt \\ e^{at} \sin bt \end{pmatrix}, \quad \mathbf{y}^2 = \begin{pmatrix} -e^{at} \sin bt \\ e^{at} \cos bt \end{pmatrix}.$$

Rinviando al paragrafo successivo per maggiori dettagli, per semplicità di notazione chiamiamo $Y(a, b)(t)$ la matrice risolvente specificata sopra, ossia

$$Y(a, b)(t) := e^{at} \begin{pmatrix} \cos bt & -\sin bt \\ \sin bt & \cos bt \end{pmatrix}. \quad (2.45)$$

Dunque, nel caso in cui A abbia gli autovalori reali (contati secondo molteplicità) $\lambda_1, \dots, \lambda_k$ e gli autovalori complessi $a_1 \pm ib_1, \dots, a_r \pm ib_r$ e inoltre $\mathbf{m}^1, \dots, \mathbf{m}^k$ ($\boldsymbol{\mu}^1 \pm i\boldsymbol{\nu}^1, \dots, \boldsymbol{\mu}^r \pm i\boldsymbol{\nu}^r$) siano gli autovettori associati ai λ_i (agli $a_j \pm ib_j$, rispettivamente), allora la matrice $M := (\mathbf{m}^1 \dots \mathbf{m}^k \boldsymbol{\nu}^1 \boldsymbol{\mu}^1 \dots \boldsymbol{\nu}^r \boldsymbol{\mu}^r)$ trasforma il sistema nella forma $\mathbf{z}' = D\mathbf{z}$, ove $D = \text{diag}(\lambda_1, \dots, \lambda_k, B_1, \dots, B_r)$ e le matrici $B_j \in \mathcal{M}(\mathbb{R}^2)$ sono come in (2.13). Dunque, la matrice risolvente nelle coordinate $\mathbf{z}$ sarà data da

$$Z(t) = \text{diag}(e^{\lambda_1 t}, \dots, e^{\lambda_k t}, Y(a_1, b_1)(t), \dots, Y(a_r, b_r)(t)), \quad (2.46)$$

da cui sarebbe possibile trarre anche una formula esplicita per una matrice fondamentale $Y = MZ$ di (2.40) sulla falsariga di (2.44), formula che evito di riportare per non appesantire i conti.

2.5 Sistemi a coefficienti costanti non diagonalizzabili

Vorremmo ripetere anche in questo caso il procedimento svolto nel paragrafo precedente. Data infatti una qualunque matrice $A \in \mathcal{M}(\mathbb{R}^N)$, una volta scritto $A = MJM^{-1}$, ove J è la forma canonica di Jordan (forma *reale*, qualora A abbia autovalori complessi) e $M \in \mathcal{M}(\mathbb{R}^N)$ la matrice (reale) di trasformazione, sappiamo che il cambiamento di coordinate $\mathbf{y} = M\mathbf{z}$ trasforma il sistema (2.40) nella forma più semplice $\mathbf{z}' = J\mathbf{z}$. Tuttavia, poiché J non è diagonale, la risoluzione del sistema trasformato non è immediata come nel caso precedente. Dovremo dunque dapprima sviluppare una formula che ci permetta di determinare un'espressione "astratta" della soluzione per qualsiasi scelta di J (anche J uguale alla A di partenza) e poi vedere come si può ricavare da questa formula una tecnica risolutiva concreta nel caso in cui J sia la forma canonica di Jordan di A . L'elemento fondamentale è la seguente

Definizione 2.26. Sia $A \in \mathcal{M}(\mathbb{C}^N)$. Poniamo allora

$$\exp(A) = e^A := \sum_{k=0}^{\infty} \frac{A^k}{k!}. \quad (2.47)$$

La matrice $\exp(A) \in \mathcal{M}(\mathbb{C}^N)$ si dice matrice esponenziale associata ad A .

La definizione data è analoga alla costruzione per serie dell'esponenziale complesso. Naturalmente, perchè abbia senso è necessario mostrare la

Proposizione 2.27. Per qualsiasi scelta di $A \in \mathcal{M}(\mathbb{C}^N)$, la serie a secondo membro di (2.47) è convergente.

Prova. In effetti, la serie converge nel senso della *convergenza totale* (cfr. Teorema 1.2). Infatti, grazie al Lemma 2.16, si ha

$$\frac{\|A^k\|_{\mathcal{L}}}{k!} \leq \frac{\|A\|_{\mathcal{L}}^k}{k!}$$

e dunque la serie in (2.47) converge e la norma $\|\cdot\|_{\mathcal{L}}$ della somma è maggiorata da $\exp(\|A\|_{\mathcal{L}})$. ■

Il nome “matrice esponenziale” è giustificato anche dalla seguente proprietà, di cui omettiamo la dimostrazione (non difficile, ma leggermente tecnica – vedi [3, p. 391]).

Lemma 2.28. *Siano $A, B \in \mathcal{M}(\mathbb{C}^N)$ tali che $AB = BA$. Allora, si ha che*

$$e^A e^B = e^{A+B}. \quad (2.48)$$

La conseguenza che ci interessa di più è data dal

Teorema 2.29. *La matrice risolvente del sistema omogeneo (2.40) è*

$$Y(t) = \exp(At). \quad (2.49)$$

Di conseguenza, la soluzione del relativo problema di Cauchy (2.15) (per $t_0 = 0$) è $\mathbf{y} = e^{At}\mathbf{y}_0$.

Prova. Mostriamo che la matrice Y definita in (2.49) risolve l’equazione matriciale $Y' = AY$. Per definizione di derivata, dobbiamo mostrare che

$$\lim_{h \rightarrow 0} \frac{Y(t+h) - Y(t)}{h} - AY = 0.$$

Si ha

$$\begin{aligned} \frac{Y(t+h) - Y(t)}{h} - AY &= \frac{e^{A(t+h)} - e^{At}}{h} - Ae^{At} \\ &= \frac{e^{At}}{h} (e^{Ah} - I - Ah) = \frac{e^{At}}{h} \left(\sum_{k=0}^{\infty} \frac{A^k h^k}{k!} - I - Ah \right), \end{aligned}$$

da cui calcolando la norma $\|\cdot\|_{\mathcal{L}}$, si ha

$$\left\| \frac{Y(t+h) - Y(t)}{h} - AY \right\|_{\mathcal{L}} \leq \|e^{At}\|_{\mathcal{L}} |h| \|A^2\|_{\mathcal{L}} \sum_{j=0}^{\infty} \frac{h^j \|A\|_{\mathcal{L}}^j}{(j+2)!}.$$

Passando al limite per $h \rightarrow 0$, si vede facilmente che il secondo membro tende a 0 e dunque altrettanto fa il primo. ■

Il teorema dimostrato peraltro di per sè dice poco di pratico: per renderlo utile, dobbiamo riuscire a calcolare e^{At} almeno per qualche scelta esplicita di A . A questo scopo, osserviamo innanzitutto che, grazie alla definizione stessa di matrice esponenziale, se $A = MJM^{-1}$ per certe $M, J \in \mathcal{M}(\mathbb{R}^N)$, allora

$$e^{At} = Me^{Jt}M^{-1};$$

dunque, è sufficiente saper calcolare e^{Jt} quando J è la forma di Jordan della matrice A . Supponiamo allora che $\lambda_1, \dots, \lambda_N$ siano gli autovalori di A contati secondo molteplicità e per il momento supposti reali. Utilizzando la decomposizione $J = D + G$, ove $D = \text{diag}(\lambda_1, \dots, \lambda_N)$ e G è come nell’Oss. 2.11, è facile vedere che

$$e^{At} = Me^{Jt}M^{-1} = Me^{(D+G)t}M^{-1} = M \text{diag}(e^{\lambda_1 t}, \dots, e^{\lambda_N t}) e^{Gt} M^{-1}.$$

Si noti infatti che $DG = GD$ grazie al Teorema 2.5.

Esempio 2.30. Consideriamo il caso in cui A è diagonalizzabile. Si ha allora $G = 0$ e $Y(t) = e^{At} = Me^{Dt}M^{-1}$. In effetti, se consideriamo la c.i. $\mathbf{y}(0) = \mathbf{y}_0$ e chiamiamo $\mathbf{y}$ la soluzione del corrispondente problema di Cauchy, abbiamo $\mathbf{y}(t) = Y(t)\mathbf{y}_0 = Me^{Dt}M^{-1}\mathbf{y}_0$. Ponendo $\mathbf{z} := M^{-1}\mathbf{y}$, $\mathbf{z}_0 := M^{-1}\mathbf{y}_0$, ritroviamo che $\mathbf{z}(t) = e^{Dt}\mathbf{z}_0$, in accordo col paragrafo precedente.

Ricordando ora che $G = \text{diag}(G_1, \dots, G_k)$, ove k è il numero degli autovalori *distinti* di A , e che ciascuna delle G_j ha la decomposizione (2.10), vediamo che basta calcolare $e^{G_s t}$ quando G_s è un singolo blocco della forma (2.9). Utilizzando l'Esercizio 2.4, determinare una formula esplicita è semplice. Non altrettanto semplice è scriverla:

$$\begin{aligned} G_s = \begin{pmatrix} \mathbf{0}^{\tau(s)} & I_{\tau(s)} \\ 0 & \mathbf{0}_{\tau(s)} \end{pmatrix} &\implies e^{G_s t} = I_{\tau(s)+1} + t \begin{pmatrix} \mathbf{0}^{\tau(s)} & I_{\tau(s)} \\ 0 & \mathbf{0}_{\tau(s)} \end{pmatrix} \\ &+ \frac{t^2}{2} \begin{pmatrix} \mathbf{0}^{\tau(s)-1} & \mathbf{0}^{\tau(s)-1} & I_{\tau(s)-1} \\ 0 & 0 & \mathbf{0}_{\tau(s)-1} \\ 0 & 0 & \mathbf{0}_{\tau(s)-1} \end{pmatrix} + \dots + \frac{t^{\tau(s)}}{\tau(s)!} U_{\tau(s)+1}, \end{aligned} \quad (2.50)$$

dove $U_{\tau(s)+1} = (u_{ij}) \in \mathcal{M}(\mathbb{R}^{\tau(s)+1})$ è la matrice avente tutti gli elementi nulli salvo $u_{1,\tau(s)+1} = 1$.

Per chiarire il significato della formula precedente, esplicitiamo l'espressione di $e^{J_2 t}$, $e^{J_3 t}$ quando $J_2 \in \mathcal{M}(\mathbb{R}^2)$, $J_3 \in \mathcal{M}(\mathbb{R}^3)$ sono date da un solo blocco di Jordan associato all'autovalore λ . In tal caso, si ha, rispettivamente,

$$e^{J_2 t} = e^{\lambda t} \begin{pmatrix} 1 & t \\ 0 & 1 \end{pmatrix}, \quad e^{J_3 t} = e^{\lambda t} \begin{pmatrix} 1 & t & t^2/2 \\ 0 & 1 & t \\ 0 & 0 & 1 \end{pmatrix}. \quad (2.51)$$

Nel caso vi siano autovalori complessi, utilizzando la forma di Jordan reale e ripetendo il ragionamento svolto si arriva ad un'espressione analoga. Ricordiamo infatti che, se B è come in (2.13), allora $e^{Bt} = Y(a, b)(t)$ (cfr. (2.45)) è la matrice risolvante del sistema $\mathbf{y}' = B\mathbf{y}$. Anzi, suggeriamo al lettore di dimostrarlo nei seguenti modi, entrambi istruttivi:

1. risolvere esplicitamente per sostituzione il sistema $\mathbf{y}' = B\mathbf{y}$;
2. Porre $C := B - aI$. Osservando che $(aI)C = C(aI)$, si può usare la proprietà (2.48). Calcolare le prime quattro potenze di C e quindi calcolare e^{Ct} utilizzando la definizione. A che cosa danno luogo le potenze pari di C ? A che cosa quelle dispari?

Dunque, nel caso vi siano autovalori complessi, l'espressione e^{Dt} sarà ancora data dalla $Z(t)$ in (2.46), mentre la e^{Gt} avrà un'espressione pressoché analoga a quella che si aveva nel caso reale.

Dopo tutte queste formule, torniamo ad un punto di vista più generale. Piuttosto che sapere risolvere esplicitamente il generico sistema $N \times N$ a coefficienti costanti (si noti peraltro che in dimensione bassa i casi tecnicamente più complicati non si presentano), è utile esplicitare la seguente conseguenza del ragionamento svolto finora:

Teorema 2.31. Sia $A \in \mathcal{M}(\mathbb{R}^N)$ e siano $\lambda_1, \dots, \lambda_k$ gli autovalori reali (rispettivamente $a_1 \pm ib_1, \dots, a_r \pm ib_r$ quelli complessi) distinti di A . Allora ogni componente di ogni soluzione del sistema (2.40) è una combinazione lineare di funzioni della forma

$$t^\alpha e^{\lambda_i t}, \quad t^\alpha e^{a_j t} \cos b_j t, \quad t^\alpha e^{a_j t} \sin b_j t, \quad (2.52)$$

ove $i = 1, \dots, k, j = 1, \dots, r, \alpha \geq 0$. ■

In sostanza, dunque, un sistema lineare omogeneo a coefficienti costanti può avere come soluzioni soltanto combinazioni lineari di prodotti di potenze per esponenziali o funzioni trigonometriche. Questa caratterizzazione delle soluzioni sarà fondamentale nella teoria della stabilità, argomento del prossimo capitolo.

Esempio 2.32. Per concludere il paragrafo, andiamo a considerare più in dettaglio il problema di Cauchy

$$\begin{cases} \mathbf{y}' = A\mathbf{y} \\ \mathbf{y}(0) = \mathbf{y}_0, \end{cases} \quad (2.53)$$

dove A è una qualunque matrice quadrata, e andiamo a chiederci che cosa succede se $\mathbf{y}_0$ è un autovettore $\mathbf{m}$ associato ad un autovalore λ di A . Si vede subito che la soluzione $\mathbf{y}$ è data da

$$\mathbf{y}(t) = e^{At} \mathbf{m} = \sum_{j=0}^{\infty} \frac{t^j A^j \mathbf{m}}{j!} = \sum_{j=0}^{\infty} \frac{t^j \lambda^j \mathbf{m}}{j!} = e^{\lambda t} \mathbf{m}; \quad (2.54)$$

dunque, se il dato iniziale è autovettore della matrice A , la corrispondente soluzione resta nello stesso autospazio per ogni tempo t .

Supponiamo invece che $\mathbf{y}_0$ sia un autovettore generalizzato di A e, per fissare le idee, immaginiamo che sia $\mathbf{y}_0 = \mathbf{m}^2$, ove $A\mathbf{m}^2 - \lambda\mathbf{m}^2 = \mathbf{m}^1$ e $A\mathbf{m}^1 - \lambda\mathbf{m}^1 = 0$. Imitando il conto (2.54), un facile procedimento per induzione (fare!) mostra che

$$\mathbf{y}(t) = \sum_{j=0}^{\infty} \frac{t^j A^j \mathbf{m}^2}{j!} = \sum_{j=0}^{\infty} \frac{t^j}{j!} (\lambda^j \mathbf{m}^2 + j \lambda^{j-1} \mathbf{m}^1) = e^{\lambda t} \mathbf{m}^2 + t e^{\lambda t} \mathbf{m}^1. \quad (2.55)$$

Dunque, anche in questo caso la traiettoria non esce dall'autospazio generalizzato.

Esercizio 2.33. Ritrovare la conclusione di (2.54) esplicitando $A = MJM^{-1}$ nello sviluppo della serie che definisce e^{At} .

Esercizio 2.34. Calcolare la soluzione del problema di Cauchy

$$\begin{cases} \mathbf{y}' = \begin{pmatrix} \lambda & 1 \\ 0 & \lambda \end{pmatrix} \mathbf{y} \\ \mathbf{y}(0) = \begin{pmatrix} 0 \\ 1 \end{pmatrix}. \end{cases} \quad (2.56)$$

2.6 Equazioni lineari a coefficienti costanti di ordine superiore al primo

In questo paragrafo ci occupiamo dapprima della determinazione dell'integrale generale dell'equazione lineare omogenea di ordine k a coefficienti costanti

$$\sum_{i=0}^k a_i y^{(i)}(t) = 0, \quad a_k \neq 0; \quad (2.57)$$

in un secondo momento, descriveremo un metodo, alternativo e più semplice rispetto a quello della variazione delle costanti, per trattare l'equazione non omogenea associata alla (2.57). Questo metodo si applicherà tuttavia soltanto per secondi membri $b(t)$ di alcune forme particolari.

Osserviamo innanzitutto che, come nel caso dei sistemi a coefficienti costanti, si applicano il Teorema 1.23 di esistenza e unicità in grande e il Teorema 2.19 di caratterizzazione delle soluzioni, che assicurano che la (2.57) ammette esattamente k soluzioni linearmente indipendenti definite su tutto $\mathbb{R}$. La seguente definizione è fondamentale:

Definizione 2.35. *L'equazione algebrica di grado k*

$$\sum_{i=0}^k a_i \lambda^i = 0 \quad (2.58)$$

viene detta equazione caratteristica associata alla (2.57).

Indicheremo nel seguito con D l'operatore di derivazione rispetto alla variabile t . Introduciamo allora l'espressione

$$P(D) = \sum_{i=0}^r b_i D^i, \quad \text{ove } b_1, \dots, b_r \in \mathbb{R}, \quad (2.59)$$

che corrisponde formalmente a un polinomio nel simbolo D . L'oggetto $P(D)$ può essere inteso anch'esso come un operatore di derivazione nel senso che, data una funzione $u \in C^r(\mathbb{R})$, si ponga:

$$P(D)u = \sum_{i=0}^r b_i u^{(i)}.$$

Naturalmente, l'operatore di derivazione $P(D)$ è lineare. In realtà, la corrispondenza tra polinomi e operatori di derivazione è più profonda. Infatti, se $P(D)$, $Q(D)$ sono operatori della forma (2.59), allora si ha

$$P(D) \circ Q(D) = (PQ)(D), \quad (2.60)$$

ossia, la composizione degli operatori differenziali P e Q si ottiene facendo il prodotto dei polinomi corrispondenti. In termini esatti, si può dire che c'è un isomorfismo tra l'anello $\mathbb{R}[D]$ dei polinomi in D a coefficienti reali dotato della somma e del prodotto usuali e l'anello degli operatori differenziali a coefficienti costanti dotato della somma

e del prodotto di composizione. Una verifica rigorosa di questo fatto richiederebbe una dimostrazione per induzione sul grado; il lettore può limitarsi a convincersene tramite qualche esempio.

Si noti che, riscrivendo la (2.57) nella forma $P(D)y = 0$, segue subito che il polinomio P è lo stesso che compare (scritto nella variabile λ) a primo membro dell'equazione caratteristica (2.58). D'altro canto, grazie al teorema fondamentale dell'algebra sappiamo che esistono numeri reali distinti $\lambda_1, \dots, \lambda_r$, numeri complessi distinti $\alpha_1 \pm i\beta_1, \dots, \alpha_s \pm i\beta_s$ e numeri naturali $m_i, n_j \geq 1$, $i = 1, \dots, r$, $j = 1, \dots, s$, tali che

$$P(\lambda) = \prod_{i=1}^r (\lambda - \lambda_i)^{m_i} \prod_{j=1}^s (\lambda - (\alpha_j + i\beta_j))^{n_j} (\lambda - (\alpha_j - i\beta_j))^{n_j}; \quad (2.61)$$

sostituendo λ con D ed interpretando $\prod$ come prodotto di composizione, otteniamo ovviamente una analoga fattorizzazione dell'operatore differenziale $P(D)$.

A questo punto, con un facile conto si dimostra che

$$\begin{aligned} (D - \lambda_i)e^{\lambda_i t} &\equiv 0, \\ (D - (\alpha_j + i\beta_j))(D - (\alpha_j - i\beta_j))e^{\alpha_j t} \cos(\beta_j t) &\equiv 0, \\ (D - (\alpha_j + i\beta_j))(D - (\alpha_j - i\beta_j))e^{\alpha_j t} \sin(\beta_j t) &\equiv 0. \end{aligned}$$

Inoltre, è ovvio che, se $Q(D)$ è un qualsiasi operatore differenziale a coefficienti costanti, allora $Q(D)0 = 0$, ossia $Q(D)$ applicato alla funzione nulla restituisce la funzione nulla.

Mettendo assieme i ragionamenti fatti fino a questo momento, saremmo a posto se non ci fosse il problema delle radici multiple. A questo proposito, dimostriamo nel caso delle radici reali una Proposizione che il lettore estenderà facilmente a quelle complesse:

Proposizione 2.36. *Sia $\lambda \in \mathbb{R}$, $m \in \mathbb{N}$, $m \geq 1$. Allora le funzioni*

$$e^{\lambda t}, \quad te^{\lambda t}, \quad \dots, \quad t^{m-1}e^{\lambda t} \quad (2.62)$$

sono linearmente indipendenti e inoltre si ha

$$(D - \lambda)^m (t^j e^{\lambda t}) \equiv 0 \quad \forall j = 0, \dots, m-1. \quad (2.63)$$

Prova. L'indipendenza lineare è banale grazie alle proprietà di polinomi ed esponenziali. Per quanto riguarda la seconda proprietà, mostriamo, per induzione su j , che $(D - \lambda)^j t^{j-1} e^{\lambda t} \equiv 0$. È ovvio che, se $j = 1$, si ha $(D - \lambda)e^{\lambda t} = 0$. Ora, supponiamo che la proprietà voluta valga per j . Si ha

$$(D - \lambda)^{j+1} (t^j e^{\lambda t}) = (D - \lambda)^j (D - \lambda)(t^j e^{\lambda t}) = (D - \lambda)^j (j t^{j-1} e^{\lambda t})$$

e concludiamo grazie all'ipotesi di induzione. ■

Ricapitolando, abbiamo quasi dimostrato il seguente

Teorema 2.37. *Se il polinomio caratteristico $P(\lambda)$ ha la decomposizione (2.61), allora una base di soluzioni per l'equazione (2.57) è data dalle funzioni*

$$\begin{aligned} e^{\lambda_i t}, \quad \dots, \quad t^{m_i-1} e^{\lambda_i t}, \quad i = 1, \dots, r, \\ e^{\alpha_j t} \cos(\beta_j t), \quad \dots, \quad t^{n_j-1} e^{\alpha_j t} \cos(\beta_j t), \quad j = 1, \dots, s, \\ e^{\alpha_j t} \sin(\beta_j t), \quad \dots, \quad t^{n_j-1} e^{\alpha_j t} \sin(\beta_j t), \quad j = 1, \dots, s. \quad \blacksquare \end{aligned}$$

Il “quasi” si riferisce al fatto che sappiamo che le funzioni elencate sopra sono in numero di k e sono tutte soluzioni di (2.57); non abbiamo però dimostrato che sono linearmente indipendenti. Tuttavia, dato che la verifica di questa naturale proprietà è piuttosto tecnica e noiosa, preferiamo soprassedere.

Concludiamo trattando il caso dell'equazione non omogenea, qualora il termine noto $b(t)$ sia di una delle seguenti forme:

$$b(t) = A(t)e^{\lambda t}, \quad b(t) = A(t) \cos(\mu t), \quad b(t) = A(t) \sin(\mu t) \quad (2.64)$$

(ove $\lambda, \mu \in \mathbb{R}$ e $A(t) \in \mathbb{R}[t]$), ovvero sia una combinazione lineare di termini di questi tre tipi. Ci limitiamo a dare una ricetta concreta per cercare una soluzione particolare dell'equazione completa, senza dimostrare nulla. In realtà le prove, tecniche ma non difficili, rispecchiano abbastanza i tipi di ragionamenti visti nel caso omogeneo.

Teorema 2.38. *Sia $b(t) = A(t)e^{\lambda t}$ e supponiamo che λ sia radice di molteplicità h (con $h \geq 0$) del polinomio caratteristico P della (2.57). Allora esiste una soluzione $\bar{y}$ dell'equazione completa*

$$\sum_{i=0}^k a_i y^{(i)}(t) = b(t) \quad (2.65)$$

della forma particolare

$$\bar{y}(t) = t^h B(t) e^{\lambda t}, \quad (2.66)$$

ove $B(t)$ è un polinomio di grado al più uguale al grado di A . Analogamente, se $b(t) = A_1(t) \cos(\mu t) + A_2(t) \sin(\mu t)$ e μi è radice di molteplicità h di P , allora esiste una soluzione $\bar{y}$ di (2.65) della forma particolare

$$\bar{y}(t) = t^h (B_1(t) \cos(\mu t) + B_2(t) \sin(\mu t)), \quad (2.67)$$

ove $B_1(t), B_2(t)$ sono polinomi di grado al più uguale al massimo dei gradi di A_1, A_2 .

Naturalmente, per linearità, qualora b sia una combinazione lineare di più termini delle forme (2.64), potremo trovare una $\bar{y}$ che sarà a sua volta una combinazione lineare di funzioni dei tipi (2.66) e (2.67).

Osservazione 2.39. Si osservi che, se il termine noto è, ad esempio, del tipo $b(t) = A(t) \cos(\mu t)$ (o l'analogo col seno al posto del coseno), la soluzione $\bar{y}$ va comunque cercata nella forma (2.67) (con sia seno che coseno).

Esercizio 2.40. Il lettore provi a congetturare che tipo di soluzione $\bar{y}$ va cercata nel caso si abbia

$$b(t) = e^{\lambda t} (A_1(t) \cos(\mu t) + A_2(t) \sin(\mu t)).$$

3 Stabilità

In questo capitolo torniamo ad occuparci del Problema di Cauchy Autonomo ((PCA) nel seguito) in avanti del primo ordine

$$\begin{cases} \mathbf{y}'(t) = \mathbf{f}(\mathbf{y}(t)) & \text{per } t \geq 0, \\ \mathbf{y}(0) = \mathbf{y}_0. \end{cases} \quad (\text{PCA})$$

Da qui in poi, supporremo sempre che Ω sia un aperto di $\mathbb{R}^N$ e che sia $\mathbf{f} \in C^1(\Omega; \mathbb{R}^N)$. Inoltre, $\mathbf{y}_0$ sarà un punto di Ω e sceglieremo come tempo iniziale $t_0 = 0$, cosa lecita poiché il sistema è autonomo. Sappiamo allora che il problema ammette un'unica soluzione massimale $(\mathbf{y}, T)$; inoltre, se $\Omega = \mathbb{R}^N$ e $\mathbf{f}$ verifica l'ipotesi di sottolinearità (1.29), possiamo concludere che $T = +\infty$; cioè, la soluzione è definita “in grande”.

In realtà, la condizione $T = +\infty$, che sappiamo in molti casi valere anche se non è verificata (1.29) avrà un'importanza centrale in quasi tutti i risultati che vedremo; anzi, spesso la possibilità di escludere il *blow up* in tempi finiti sarà una conseguenza delle nuove ipotesi. In effetti, nel corso del capitolo ci occuperemo prevalentemente del caso in cui $T = +\infty$ e affronteremo le seguenti questioni:

1. Esiste il limite $\ell = \lim_{t \rightarrow +\infty} \mathbf{y}(t)$? Se esiste, posso dare delle condizioni sotto le quali ℓ è finito?
2. È possibile dare una caratterizzazione dei possibili limiti ℓ al variare del dato iniziale $\mathbf{y}_0$?
3. Nel caso in cui $\ell \in \mathbb{R}^N$, posso stimare la velocità di convergenza? Se invece $\ell = \infty$ (ossia $|\mathbf{y}(t)| \rightarrow +\infty$), posso stimare la velocità di divergenza?
4. Infine, nei casi in cui il limite ℓ non esiste, cosa può succedere alla traiettoria $\mathbf{y}(t)$ quando t “diventa grande”?

Nel caso lineare, quando sicuramente $T = +\infty$, potremo dare una risposta completa a queste domande utilizzando la caratterizzazione delle soluzioni fornita nel capitolo precedente (Teorema 2.31). In particolare, potremo individuare i concetti fondamentali di “punto di equilibrio” e di “stabilità”. Questi saranno ripresi e ulteriormente dettagliati nello studio del caso non lineare, dove la situazione è molto più complessa e saremo in grado di ottenere solamente risposte parziali alle domande 1.–4. .

3.1 Stabilità dei sistemi lineari

Sia $A \in \mathcal{M}(\mathbb{R}^N)$. Consideriamo le soluzioni del Problema di Cauchy Autonomo Lineare ((PCAL) nel seguito) in avanti

$$\begin{cases} \mathbf{y}'(t) = A\mathbf{y}(t) & \text{per } t \geq 0, \\ \mathbf{y}(0) = \mathbf{a} \end{cases} \quad (\text{PCAL})$$

al variare del dato iniziale $\mathbf{a} \in \mathbb{R}^N$. Qui, evidentemente, $\Omega = \mathbb{R}^N$ e $A \in \mathcal{M}(\mathbb{R}^N)$; sto dunque considerando il caso omogeneo e a coefficienti costanti (altrimenti il sistema non sarebbe autonomo). Presentiamo innanzitutto alcune notazioni, molte delle quali saranno usate anche per il caso non lineare.

Già sappiamo che la soluzione di (PCAL) può essere interpretata come la *traiettoria* di un punto materiale che parte all'istante $t = 0$ dalla posizione $\mathbf{a}$. Per questo motivo, sarà conveniente indicarla con la notazione $\mathbf{y}_{\mathbf{a}}(t)$. Naturalmente, $\mathbf{y}_{\mathbf{a}}$ avrà tempo finale $T = +\infty$. Vale la pena inoltre ricordare (vedi il Paragrafo 1.6) che, se $\mathbf{y}_{\mathbf{a}}, \mathbf{y}_{\mathbf{b}}$ sono due soluzioni relative alle posizioni iniziali $\mathbf{a}, \mathbf{b} \in \mathbb{R}^N$, allora

$$\mathbf{y}_{\mathbf{a}}(t_1) = \mathbf{y}_{\mathbf{b}}(t_2) \implies \mathbf{y}_{\mathbf{a}}(t_1 + t) = \mathbf{y}_{\mathbf{b}}(t_2 + t) \quad \forall t \in \mathbb{R}; \quad (3.1)$$

in particolare, si ha

$$\mathbf{y}_a(t_1 - t_2) = \mathbf{y}_b(0) = \mathbf{b}, \quad \mathbf{y}_b(t_2 - t_1) = \mathbf{y}_a(0) = \mathbf{a}$$

(nel caso non lineare queste relazioni continuano a valere per scelte *ammissibili* di t).

Sia ora data una traiettoria $\mathbf{y}_a$ del sistema (PCA) (consideriamo dunque il caso generale), supponiamo che abbia tempo finale $+\infty$ e che tenda a $\ell \in \Omega$ per $t \rightarrow +\infty$. Dalla continuità di $\mathbf{f}$ segue allora

$$\lim_{t \rightarrow +\infty} \mathbf{y}'_a(t) = \lim_{t \rightarrow +\infty} \mathbf{f}(\mathbf{y}_a(t)) = \mathbf{f}(\ell),$$

D'altronde, per il Teorema dell'asintoto, deve essere $\lim_{t \rightarrow +\infty} \mathbf{y}'_a(t) = \mathbf{0}$, e quindi $\mathbf{f}(\ell) = \mathbf{0}$. Ne segue che i limiti in Ω delle traiettorie di (PCA) (e, in particolare, di (PCAL)) sono piuttosto particolari:

Definizione 3.1. *Dico che $\bar{\mathbf{y}} \in \Omega$ è un punto di equilibrio o, anche, un punto critico per il sistema (PCA) (ovvero (PCAL)) qualora $\mathbf{f}(\bar{\mathbf{y}}) = \mathbf{0}$ (rispettivamente $A\bar{\mathbf{y}} = \mathbf{0}$).*

Si noti che nel caso non lineare (PCA), ove in generale $\Omega \neq \mathbb{R}^N$, può capitare che la traiettoria tenda per $t \rightarrow T^-$ (tempo finale, ora non necessariamente pari a $+\infty$) a un limite finito ℓ , ma con $\ell \notin \Omega$; in questo caso ovviamente ℓ non è un punto di equilibrio.

Gli zeri della funzione $\mathbf{f}$, oltre a dare luogo a traiettorie costanti, hanno quindi un interesse particolare: sono gli unici punti del dominio di $\mathbf{f}$ a cui le generiche traiettorie di (PCA) possono avvicinarsi per tempi lunghi. Naturalmente, non sempre questo avviene ed il nostro prossimo obiettivo sarà proprio quello di discutere la casistica che può presentarsi, cominciando dal caso lineare.

La prima distinzione da fare per il sistema (PCAL) è la seguente: o la matrice A ha rango N , e allora $\mathbf{0}$ è l'unico punto di equilibrio, oppure esiste almeno un vettore $\mathbf{v}$ non nullo tale che $A\mathbf{v} = \mathbf{0}$. L'insieme dei punti di equilibrio, ossia il nucleo di A , risulta allora essere una varietà lineare non banale. Questa è una situazione in un certo senso degenerare, perché i punti critici del sistema non sono isolati: la discuteremo brevemente a fine paragrafo. Il caso di gran lunga più rilevante è invece quello in cui A ha rango N , su cui ora concentriamo la nostra attenzione. Per affrontarlo, ci serve innanzitutto un lemma di algebra lineare:

Lemma 3.2. *Sia $A \in \mathcal{M}(\mathbb{R}^N)$ e supponiamo che ogni autovalore di A abbia parte reale strettamente negativa. Allora esistono un prodotto scalare $((\cdot, \cdot))$ su $\mathbb{R}^N$ ed una costante $b > 0$ tali che*

$$((A\xi, \xi)) \leq -b\|\xi\|^2 \quad \forall \xi \in \mathbb{R}^N, \quad (3.2)$$

ove $\|\cdot\|$ è la norma associata a tale prodotto scalare. Analogamente, se ogni autovalore di A ha parte reale strettamente positiva, si ha che

$$((A\xi, \xi)) \geq b\|\xi\|^2 \quad \forall \xi \in \mathbb{R}^N. \quad (3.3)$$

Prova. Diamo la dimostrazione solo in un caso molto particolare, quello cioè in cui A è diagonale e ha solo autovalori reali. Supponiamo che questi siano tutti negativi; nel caso siano positivi la dimostrazione è analoga. Scelto allora $\beta > 0$ tale che $\lambda_i \leq -\beta$ per ogni autovalore λ_i , notiamo che

$$(A\xi, \xi) = \left(A \sum_i \xi_i e^i, \sum_j \xi_j e^j \right) = \left(\sum_i \lambda_i \xi_i e^i, \sum_j \xi_j e^j \right) = \sum_{i,j} \lambda_i \xi_i \xi_j (e^i, e^j); \quad (3.4)$$

dunque vale la (3.2) rispetto al prodotto scalare euclideo e con $b = \beta$. Lo stesso conto si adatta con modifiche ovvie al caso in cui $A = \text{diag}(\lambda_1, \dots, \lambda_r, B_1, \dots, B_s)$, ove le B_j sono come in (2.13).

Nel caso generale, si procede attraverso cambiamenti di coordinate; se A è diagonalizzabile si passa alla forma diagonale; il cambiamento di coordinate individuerà il nuovo prodotto scalare rispetto al quale vale la relazione (3.2). Ancora più macchinoso è il caso in cui A non è diagonalizzabile; infatti la forma di Jordan che abbiamo introdotto nello scorso capitolo non è quella giusta per mostrare un analogo della (3.2), e dunque va (leggermente) modificata. Il lettore interessato può trovare i dettagli tecnici ad esempio su [5, p. 148]. ■

Presentiamo ora il risultato fondamentale per i sistemi lineari:

Teorema 3.3. *Sia $A \in \mathcal{M}(\mathbb{R}^N)$ invertibile. Allora le seguenti condizioni sono equivalenti:*

- (a) *Esiste $\beta > 0$ tale che $\text{Re}(\sigma) \leq -\beta$ per ogni autovalore σ di A ;*
- (b) *Esistono $b, c > 0$ tali che, per ogni $\mathbf{a} \in \mathbb{R}^N$, si ha*

$$|\mathbf{y}_{\mathbf{a}}(t)| \leq c|\mathbf{a}|e^{-bt} \quad \forall t \geq 0; \quad (3.5)$$

- (c) *Esistono $b, c > 0$ tali che $\|e^{At}\|_{\mathcal{L}} \leq ce^{-bt}$ per ogni $t \geq 0$.*

Prova. (c) $\Rightarrow$ (b). Basta notare che $\mathbf{y}_{\mathbf{a}}(t) = e^{At}\mathbf{a}$, prendere il modulo e usare (2.20).

(b) $\Rightarrow$ (c). Dividere (3.5) per $|\mathbf{a}|$ e passare al sup rispetto ad $\mathbf{a}$.

(b) $\Rightarrow$ (a). Sia $\lambda + i\mu$ un autovalore di A , con $\mu \geq 0$. Supponiamo per assurdo $\lambda \geq 0$. Allora, scrivendo $A = MJM^{-1}$ ove J è la forma canonica reale, sappiamo che nelle nuove coordinate $\mathbf{z} = M^{-1}\mathbf{y}$ esiste una soluzione

$$\bar{\mathbf{z}}(t) = (e^{\lambda t} \cos \mu t, e^{\lambda t} \sin \mu t, 0, \dots, 0).$$

Ponendo $\bar{\mathbf{y}} := M\bar{\mathbf{z}}$, è evidente che $|\bar{\mathbf{y}}(t)| \not\rightarrow 0$ per $t \rightarrow +\infty$. D'altro canto tale traiettoria deve originare da un qualche punto $\mathbf{a} \neq \mathbf{0}$, ossia $\bar{\mathbf{y}}(t) = e^{At}\mathbf{a} = \mathbf{y}_{\mathbf{a}}(t)$. A questo punto, la (b) fornisce subito l'assurdo.

(a) $\Rightarrow$ (b). Sia $\mathbf{a} \neq \mathbf{0}$ e consideriamo la corrispondente soluzione $\mathbf{y}_{\mathbf{a}} \neq \mathbf{0}$. Sia inoltre $((\cdot, \cdot))$ un prodotto scalare tale che vale la (3.2) fornita dal Lemma 3.2 e sia $\|\cdot\|$ la norma associata.

Poichè $\mathbf{y}_{\mathbf{a}}$ non può mai annullarsi, un semplice calcolo ci dice che

$$\frac{d}{dt} \|\mathbf{y}_{\mathbf{a}}\| = \frac{((\mathbf{y}'_{\mathbf{a}}, \mathbf{y}_{\mathbf{a}}))}{\|\mathbf{y}_{\mathbf{a}}\|} = \frac{((A\mathbf{y}_{\mathbf{a}}, \mathbf{y}_{\mathbf{a}}))}{\|\mathbf{y}_{\mathbf{a}}\|} \leq -b\|\mathbf{y}_{\mathbf{a}}\| \quad (3.6)$$

grazie al Lemma citato. Ora, dalla relazione precedente segue subito

$$\frac{d}{dt} \log \|\mathbf{y}_a(t)\| \leq -b, \quad \text{da cui, integrando,} \quad \|\mathbf{y}_a(t)\| \leq \|\mathbf{a}\| e^{-bt}.$$

A questo punto la tesi segue dall'equivalenza di norme in dimensione finita (Teorema 2.17). ■

Dal momento che il caso trattato nel risultato precedente è di notevole importanza, merita che gli diamo un nome e questo è il fine della prossima definizione. Avvisiamo tuttavia il lettore che la terminologia adottata è tutto fuorché standard:

Definizione 3.4. *Qualora una traiettoria $\mathbf{y}_a$ del sistema (PCAL) verifichi (3.5) per qualche $b, c > 0$, si dice che $\mathbf{y}_a$ decade esponenzialmente. Se valgono le condizioni del Teorema 3.3, dico che il sistema (PCAL) decade esponenzialmente. Analogamente, dico che $\mathbf{y}_a$ cresce esponenzialmente se esistono $b, c > 0$ tali che*

$$|\mathbf{y}_a(t)| \geq c|\mathbf{a}|e^{bt} \quad \forall t \geq 0. \quad (3.7)$$

Se ciò vale per ogni traiettoria $\mathbf{y}_a$, con b, c indipendenti da $\mathbf{a}$, dico che (PCAL) cresce esponenzialmente.

Osservazione 3.5. Ricalcando la dimostrazione del Teorema 3.3, è facile vedere che il sistema (PCAL) cresce esponenzialmente se e solo se esiste $\beta > 0$ tale che $\operatorname{Re}(\sigma) \geq \beta$ per ogni autovalore σ di A . In altre parole, condizioni analoghe alla (a) e alla (b) del Teorema sono equivalenti anche in questo caso. Invece, l'analogia di (c) è adesso una condizione più debole, che è verificata allorché A ha *almeno* un autovalore con parte reale positiva. In effetti, questa vale quando esiste *almeno* una traiettoria che cresce esponenzialmente. Sia infatti, ad esempio,

$$A = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}; \quad \text{allora} \quad \|e^{At}\|_{\mathcal{L}} \geq \left| \begin{pmatrix} e^t & 0 \\ 0 & e^{-t} \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right| = e^t.$$

Il sistema associato ad A possiede dunque traiettorie che crescono esponenzialmente, ma anche traiettorie che esponenzialmente decadono.

Il Teorema 3.3, e la discussione successiva, suggeriscono che il comportamento delle soluzioni di (PCAL) è strettamente legato agli autovalori della matrice A , ma naturalmente non coprono la totalità dei casi possibili. Infatti, risulta naturale chiedersi che cosa succede quando ci sono autovalori con parti reali di segno opposto o, anche, con parte reale nulla. Cominciamo ad enunciare, e dimostrare, un

Teorema 3.6. *Supponiamo che tutti gli autovalori di A abbiano parte reale diversa da 0. Allora esistono e sono unici due sottospazi E^s e E^i di $\mathbb{R}^N$, A -invarianti e non banali, tali che $\mathbb{R}^N = E^s \oplus E^i$; inoltre, esistono $b, c > 0$ tali che (3.5) vale per ogni $\mathbf{a} \in E^s$ e (3.7) vale per ogni $\mathbf{a} \in E^i$. E^s e E^i si dicono rispettivamente spazio stabile e spazio instabile di A .*

Prova. Consideriamo la forma canonica di Jordan reale J della matrice A . Ordiniamo i blocchi di J disponendo dapprima quelli relativi agli autovalori con parte

reale negativa e poi gli altri. Guardando la matrice M di cambiamento di coordinate, otteniamo allora una decomposizione

$$E^s := \text{span}\{\mathbf{m}^1, \dots, \mathbf{m}^k\}, \quad E^i := \text{span}\{\mathbf{m}^{k+1}, \dots, \mathbf{m}^N\},$$

ove $\mathbf{m}^1, \dots, \mathbf{m}^k$ sono gli autovettori, eventualmente generalizzati, associati agli autovalori con parte reale negativa e $\mathbf{m}^{k+1}, \dots, \mathbf{m}^N$ quelli associati agli autovalori con parte reale positiva. È facile vedere che tale decomposizione verifica le ipotesi richieste.

Mostriamo ora l'unicità di E^s, E^i , supponendo per assurdo che $\mathbb{R}^N = F^s \oplus F^i$ per un'altra coppia di spazi F^s e F^i con le stesse proprietà dei precedenti. In particolare, faccio vedere che $F^s \subset E^s$; l'inclusione inversa si dimostrerà allo stesso modo. Sia $\mathbf{x} \in F^s$. Grazie alla prima parte dell'enunciato, posso scrivere $\mathbf{x} = \mathbf{x}_s + \mathbf{x}_i$, con $\mathbf{x}_s \in E^s, \mathbf{x}_i \in E^i$. Ho allora

$$|e^{At}\mathbf{x}_i| = |e^{At}\mathbf{x} - e^{At}\mathbf{x}_s| \leq ce^{-bt}(|\mathbf{x}| + |\mathbf{x}_s|),$$

poichè $\mathbf{x} \in F^s$ e $\mathbf{x}_s \in E^s$, che sono spazi stabili. Segue dunque $\mathbf{x}_i = \mathbf{0}$ e la tesi. ■

Naturalmente la decomposizione data dal Teorema sussiste anche quando, ad esempio, il sistema (PCAL) decade esponenzialmente. In questo caso, tuttavia, $E^s = \mathbb{R}^N$, e $E^i = \{0\}$ è lo spazio banale.

Osservazione 3.7. Si noti che, se si considerano le coordinate $\mathbf{z}$ date da $\mathbf{z} = M^{-1}\mathbf{y}$, allora, in tali coordinate ed ordinando i blocchi di J come nella prova del precedente Teorema, lo spazio E^s risulta generato dai primi k vettori della base canonica.

Denoteremo d'ora in poi come $\mathcal{M}_0$ lo spazio delle matrici che hanno almeno un autovalore con parte reale nulla. Vediamo che cosa succede quando $A \in \mathcal{M}_0$. Raffinando leggermente il precedente Teorema (e la sua dimostrazione), $\mathbb{R}^N$ si decompone in questo caso nei sottospazi $E^s \oplus E^i \oplus E^c$, ove il nuovo sottospazio E^c generato dagli autovettori (eventualmente generalizzati) relativi agli autovalori con parte reale nulla, viene detto *spazio centro* del sistema. Oltre al caso in cui A non è invertibile (che abbiamo "scartato" all'inizio), E^c è non banale quando ci sono autovalori immaginari puri.

Il fatto che gli spazi E^s, E^i, E^c sono A -invarianti è molto importante. Sia infatti $\mathbf{a} \in \mathbb{R}^N \setminus \{0\}$ un dato iniziale. Decomponendo $\mathbf{a} = \mathbf{a}_s + \mathbf{a}_i + \mathbf{a}_c$, abbiamo che

$$\mathbf{y}_\mathbf{a}(t) = e^{At}(\mathbf{a}_s + \mathbf{a}_i + \mathbf{a}_c) = e^{At}\mathbf{a}_s + e^{At}\mathbf{a}_i + e^{At}\mathbf{a}_c,$$

ove $e^{At}\mathbf{a}_s \in E^s$, $e^{At}\mathbf{a}_i \in E^i$ e $e^{At}\mathbf{a}_c \in E^c$ per ogni t ; infatti, è facile vedere che $e^{At}E^s \subset E^s$, ecc. (vedi l'Esempio 2.32). In particolare, se il dato iniziale $\mathbf{a}$ sta, ad esempio, nello spazio stabile (ossia, $\mathbf{a}_i = \mathbf{a}_c = \mathbf{0}$), la traiettoria non esce mai da questo.

Esempio 3.8. Sia dato il problema di Cauchy

$$\mathbf{y}'(t) = \begin{pmatrix} 2 & 4 \\ 3 & 3 \end{pmatrix} \mathbf{y}, \quad \mathbf{y}(0) = \mathbf{y}_0 = \begin{pmatrix} 8 \\ 6 \end{pmatrix}.$$

La matrice A dei coefficienti è diagonalizzabile e si ha

$$A = MDM^{-1}, \quad D = \begin{pmatrix} -1 & 0 \\ 0 & 6 \end{pmatrix}, \quad M = \begin{pmatrix} 4 & 1 \\ 3 & 1 \end{pmatrix}.$$

Si vede allora che $\mathbf{y}_0$ è un autovettore associato a $\lambda = -1$. Ponendo $\mathbf{z} = M^{-1}\mathbf{y}$, si possono calcolare

$$M^{-1} = \begin{pmatrix} 1 & -1 \\ -3 & 4 \end{pmatrix}, \quad \mathbf{z}_0 = \begin{pmatrix} 2 \\ 0 \end{pmatrix}, \quad \mathbf{y} = Me^{Dt}\mathbf{z}_0 = M \begin{pmatrix} 2e^{-t} \\ 0 \end{pmatrix};$$

dunque, come ci aspettavamo, $\mathbf{y}_0$ è nello spazio stabile del sistema e la traiettoria che ha origine in $\mathbf{y}_0$ tende a $\mathbf{0}$ esponenzialmente.

Piccole perturbazioni e genericità. Vediamo ora che cosa può capitare quando prendiamo il sistema (PCAL) e ne cambiamo “di poco” i coefficienti. Notiamo peraltro che il discorso che faremo avrà una notevole importanza anche nello studio del sistema non lineare (PCA). Per prima cosa dobbiamo chiarire in termini rigorosi che cosa vuol dire questo “di poco”. L’approccio più naturale è quello di partire dalla matrice A , fissare $\varepsilon > 0$ e considerare le perturbazioni A_ε che distano da A meno di $\varepsilon > 0$ nella norma naturale $\|\cdot\|_{\mathcal{L}}$. Cominciamo subito col dare una

Definizione 3.9. Diciamo che una proprietà (P) relativa alle matrici $A \in \mathcal{M}(\mathbb{R}^N)$ (o, se si preferisce, del sistema (PCAL) ad esse associato) è generica se valgono le seguenti due condizioni:

- (i) per ogni A che verifica (P) , esiste $\varepsilon > 0$ tale che ogni $A_\varepsilon \in \mathcal{M}(\mathbb{R}^N)$ tale che $\|A - A_\varepsilon\|_{\mathcal{L}} \leq \varepsilon$ verifica ancora (P) ;
- (ii) per ogni A che non verifica (P) e per ogni $\varepsilon > 0$ esiste $A_\varepsilon \in \mathcal{M}(\mathbb{R}^N)$ tale che $\|A - A_\varepsilon\|_{\mathcal{L}} \leq \varepsilon$ e inoltre A_ε verifica (P) .

In altre parole, la (i) dice che la famiglia delle matrici che verificano (P) è aperta, la (ii) che tale famiglia è densa.

Chiamando sistema “di partenza” il sistema (PCAL) con matrice A e sistema “perturbato” quello con matrice A_ε , ha interesse anche il caso in cui vale la proprietà (i), ma non necessariamente la (ii); in una tale situazione diremo a volte che la proprietà (P) non è sensibile alle piccole perturbazioni. Viceversa, una proprietà (P) è sensibile alle piccole perturbazioni quando (P) non vale, ossia

(non i) esiste A che verifica (P) e tale che per ogni $\varepsilon > 0$ esiste $A_\varepsilon \in \mathcal{M}(\mathbb{R}^N)$ tale che $\|A - A_\varepsilon\|_{\mathcal{L}} \leq \varepsilon$ e A_ε non verifica (P) .

Cominciamo ora a vedere casi concreti di proprietà generiche e di proprietà che verificano almeno (i), osservando già da subito che ulteriori dettagli saranno dati a fine paragrafo per il caso bidimensionale. Cominciamo con un

Teorema 3.10. Sia $\mathcal{M}_{\text{dist}}$ il sottospazio di $\mathcal{M}(\mathbb{R}^N)$ costituito dalle matrici con autovalori distinti. Allora, $\mathcal{M}_{\text{dist}}$ è denso in $\mathcal{M}(\mathbb{R}^N)$.

Prova. Sia $A \in \mathcal{M}(\mathbb{R}^N)$ e sia $\varepsilon > 0$. Cerchiamo una matrice $A_\varepsilon \in \mathcal{M}_{\text{dist}}$ tale che

$$\|A - A_\varepsilon\|_{\mathcal{L}} \leq \varepsilon. \quad (3.8)$$

Grazie a un cambiamento di coordinate, assumiamo che A sia in forma canonica di Jordan e, inoltre, per semplicità supponiamo che abbia solo autovalori reali $\lambda_1, \dots, \lambda_N$,

con possibili ripetizioni. Se $A = D + Z$ è l'usuale decomposizione di A nelle parti diagonale e nilpotente, possiamo scegliere N numeri reali distinti $\lambda'_1, \dots, \lambda'_N$ tali che $|\lambda_i - \lambda'_i| \leq \varepsilon$ per $i = 1, \dots, N$. Ponendo $D' := \text{diag}(\lambda'_1, \dots, \lambda'_N)$, è allora facile vedere che $A_\varepsilon := D' + Z$ è la matrice cercata (in realtà avremo $c\varepsilon$ a secondo membro di (3.8), con c costante puramente dimensionale). ■

Esercizio 3.11. Provare ad estendere la dimostrazione precedente considerando una matrice A che abbia anche autovalori complessi.

Si noti che questo risultato implica che l'insieme $\mathcal{M}_{\text{diag}}$ delle matrici diagonalizzabili è denso in $\mathcal{M}(\mathbb{R}^N)$ (infatti contiene strettamente $\mathcal{M}_{\text{dist}}$); dunque, la proprietà (P) di diagonalizzabilità verifica la (ii) della Def. 3.9. Inoltre, modificando leggermente la dimostrazione possiamo provare il

Corollario 3.12. *L'insieme delle matrici con tutti gli autovalori diversi da 0 è denso in $\mathcal{M}(\mathbb{R}^N)$. Lo stesso vale per l'insieme $\mathcal{M}(\mathbb{R}^N) \setminus \mathcal{M}_0$ delle matrici i cui autovalori hanno tutti parte reale diversa da 0.*

Osservazione 3.13. Val la pena invece osservare che l'insieme delle matrici i cui autovalori hanno tutti parte *immaginaria* diversa da 0 non è denso in $\mathcal{M}(\mathbb{R}^N)$. Sia infatti ad esempio data la matrice

$$A = \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}$$

e supponiamo che esista una successione $\{A_n\}$ di matrici con autovalori complessi coniugati $\lambda_n \pm i\mu_n$ convergente ad A nella norma $\|\cdot\|_{\mathcal{L}}$. Poiché i coefficienti, e dunque le radici, del polinomio caratteristico sono funzioni continue dei coefficienti della matrice, seguirebbe allora che gli autovalori $\lambda_n \pm i\mu_n$ dovrebbero convergere rispettivamente a 1 e a 2; dunque, in particolare, la successione λ_n delle parti reali convergerebbe contemporaneamente a 1 e a 2, il che è impossibile.

Osservazione 3.14. Con un procedimento leggermente più laborioso si può mostrare che $\mathcal{M}_{\text{dist}}$ è anche aperto in $\mathcal{M}(\mathbb{R}^N)$ (invece $\mathcal{M}_{\text{diag}}$ *non* è aperto in $\mathcal{M}(\mathbb{R}^N)$). Dunque, se $A \in \mathcal{M}_{\text{dist}}$, siamo certi che almeno un intorno di A è ancora contenuto in $\mathcal{M}_{\text{dist}}$. Viceversa, qualora $A \in \mathcal{M}(\mathbb{R}^N) \setminus \mathcal{M}_{\text{dist}}$, in ogni intorno di A possiamo trovare un elemento di $\mathcal{M}_{\text{dist}}$. Analogamente, anche $\mathcal{M}(\mathbb{R}^N) \setminus \mathcal{M}_0$ è aperto in $\mathcal{M}(\mathbb{R}^N)$.

Vediamo ora di interpretare il precedente discorso dal punto di vista delle applicazioni. Se il sistema dinamico (PCAL) descrive una qualche situazione fisica, in molti casi in cui sussiste un'incertezza nella misura dei coefficienti della matrice A è lecito supporre “a vuoto” che questa verifichi una qualche proprietà generica (P) . In effetti, le matrici che non verificano (P) sono “poche” (perché vale (ii)); inoltre, se una matrice verifica (P) e noi ne perturbiamo “sufficientemente poco” i coefficienti, (P) , grazie a (i) , sarà ancora valida.

Come si è visto, esempi di proprietà generiche sono dunque il fatto che $A \in \mathcal{M}_{\text{dist}}$ ovvero che $A \in \mathcal{M}(\mathbb{R}^N) \setminus \mathcal{M}_0$. Guardando la seconda di queste, si vede in particolare che i casi in cui A non è invertibile sono “pochi”, ed è in una certa misura lecito

considerarli degeneri. In effetti, partendo da un sistema che ha un autovalore nullo (o, più in generale, con parte reale nulla), una piccolissima perturbazione dei coefficienti può farci uscire da $\mathcal{M}_0$ e “distruggere” lo spazio centro E^c .

Osservazione 3.15. Tuttavia, il precedente discorso va preso con le dovute attenzioni, perché c'è un altro fattore da considerare. Infatti, l'esistenza di stati stazionari non banali (ossia di autovalori nulli) potrebbe essere garantita da qualche principio fisico (ad esempio la conservazione di un'energia); in questi casi, ogni perturbazione dei coefficienti che li faccia sparire è priva di significato.

Sistemi bidimensionali. Per chiarire concretamente la teoria svolta fin qui, concludiamo il paragrafo concentrandoci sul caso particolare in cui $A \in \mathcal{M}(\mathbb{R}^2)$ (chiameremo allora x, y le coordinate scalari). Qui le cose sono un po' più semplici che nel caso generale; inoltre, sussiste una nomenclatura che è di uso ormai abbastanza corrente (e che solo in parte può essere estesa a una dimensione superiore). Notiamo che questa parte della teoria si trova anche su [9, § 4.2]. Detti λ_1, λ_2 gli autovalori di A (eventualmente con $\lambda_1 = \lambda_2$), veniamo dunque ad analizzare i vari casi che si possono presentare, continuando a supporre che la matrice A sia invertibile. In questo caso, $\mathbf{0}$ è l'unico punto di equilibrio, e può essere dei seguenti tipi:

– **Nodo (instabile):** $0 < \lambda_1 < \lambda_2$. In particolare, il sistema cresce esponenzialmente ($E^i = \mathbb{R}^2$). Questo caso non è sensibile alle “piccole” perturbazioni dei coefficienti.

Esempio 3.16. Consideriamo il sistema

$$\mathbf{y}' = \begin{pmatrix} 1 & 0 \\ 0 & 3 \end{pmatrix} \mathbf{y}.$$

Si può avere un'idea dell'andamento qualitativo delle traiettorie osservando che l'integrale generale è $\mathbf{y}_a = (a_1 e^t, a_2 e^{3t})$. Dunque è possibile eliminare la t , ottenendo ad esempio $y = a_2(x/a_1)^3$.

Quindi, se la matrice A è già in forma diagonale, le traiettorie hanno l'andamento di rami di funzioni di tipo potenza e vengono percorse in direzione uscente dall'origine; altrimenti, questo varrà rispetto alle coordinate $\mathbf{z}$. A volte questo caso, e il successivo, vengono anche detti “**nodo proprio**” o “**nodo a due tangenti**”, le tangenti essendo date dalle direzioni individuate dagli autovettori di A (che, come abbiamo visto, corrispondono alle traiettorie rettilinee).

– **Nodo (stabile):** $\lambda_1 < \lambda_2 < 0$. Questo caso è analogo al precedente, ma il verso di percorrenza delle traiettorie è invertito e il sistema decade esponenzialmente.

– **Sella:** $\lambda_1 < 0 < \lambda_2$. Lo spazio stabile e lo spazio instabile verranno ad avere entrambi dimensione 1 ed anche questo caso, almeno in dimensione 2, non è sensibile rispetto a piccole perturbazioni dei coefficienti. Passando alle coordinate $\mathbf{z}$ ed eliminando la variabile t si può vedere che le traiettorie hanno l'andamento di potenze ad esponente negativo. Notiamo peraltro che, anche in dimensione superiore, si dice a volte che $\mathbf{0}$ è punto di “sella” quando ci siano contemporaneamente almeno un autovalore con parte reale positiva e uno con parte reale negativa.

– **Nodo improprio:** $0 < \lambda_1 = \lambda_2$ (si tratta in modo analogo il caso di autovalori coincidenti e strettamente negativi). Chiaramente, esistono due sottocasi: il primo è quello in cui A è diagonalizzabile (che si chiama a volte **nodo a stella**). Allora, A è già in forma diagonale (vedi Es. 2.3). Questo vale a dire che le traiettorie vivono su semirette uscenti dall'origine.

Se invece A non è diagonalizzabile, allora si parla di **nodo a una tangente** (l'autospazio ha dimensione 1 e dunque c'è una sola traiettoria rettilinea). Il lettore può chiarirsi ulteriormente le idee andando a studiare il sistema

$$\mathbf{y}' = \begin{pmatrix} \lambda & 1 \\ 0 & \lambda \end{pmatrix} \mathbf{y}, \quad \lambda > 0,$$

che corrisponde al caso in cui A è già scritta nella forma canonica di Jordan. Vedi anche [9, Fig. 2 b), p. 256].

Sia nel caso del nodo a stella che di quello a una tangente, il sistema è “sensibile” rispetto a piccole perturbazioni dei coefficienti (queste considerazioni potrebbero estendersi a dimensioni superiori). Infatti, dato che l'equazione caratteristica ha una radice doppia, e dunque ha discriminante nullo, cambiando anche di pochissimo i coefficienti, possono capitare tre cose: o si resta nella stessa situazione (autovalore reale doppio), oppure il discriminante diventa negativo e si passa ad autovalori complessi coniugati (vedi sotto) o, infine, il discriminante diviene strettamente positivo (e allora si passa ad autovalori reali distinti, nel caso considerato strettamente positivi).

– **Fuoco (stabile o instabile):** $\lambda_{1,2} = \alpha \pm i\beta$, con $\alpha, \beta \neq 0$. Ovviamente parliamo di fuoco stabile se $\alpha < 0$ e instabile se $\alpha > 0$. Nel primo caso, il sistema cresce, nel secondo decade esponenzialmente. Nelle coordinate $\mathbf{z}$ si vede facilmente che le traiettorie descrivono spirali, nel primo caso entranti e nel secondo uscenti dall'origine. Guardando l'equazione caratteristica, questa risulta avere discriminante negativo, il che permane qualora si operi una piccola perturbazione dei coefficienti (sistema non sensibile alle piccole perturbazioni).

– **Centro:** $\lambda_{1,2} = \pm i\beta$, $\beta \neq 0$. In quest'ultimo caso, abbiamo che E^c è non banale, poiché gli autovalori hanno parte reale nulla. Il nome **centro** è ulteriormente giustificato se si disegnano le traiettorie, che risultano essere ellissi centrate nell'origine. In questa situazione una piccola perturbazione dei coefficienti non può cambiare il segno del discriminante; dunque gli autovalori resteranno complessi coniugati. Invece, può far diventare strettamente negativa o positiva la parte reale degli autovalori, cioè dar luogo a un fuoco, stabile o instabile rispettivamente. Globalmente, il sistema è dunque “sensibile” rispetto a perturbazioni.

Come abbiamo osservato all'inizio, in certi casi la terminologia che abbiamo introdotto si può mantenere anche in dimensione maggiore di 2, come nel seguente

Esempio 3.17. Consideriamo il sistema

$$\mathbf{y}' = \begin{pmatrix} 2 & 0 & 0 \\ 0 & 1 & 4 \\ 0 & 1 & 1 \end{pmatrix} \mathbf{y}$$

e proviamo a determinare gli spazi E^s , E^i , E^c . Diagonalizzando la matrice A dei coefficienti, troviamo che gli autovalori sono $\lambda_1 = 2$, $\lambda_2 = 3$, $\lambda_3 = -1$, cui corrispondono, rispettivamente, gli autovettori $\mathbf{m}^1 = (1, 0, 0)$, $\mathbf{m}^2 = (0, 2, 1)$, $\mathbf{m}^3 = (0, 2, -1)$. Dunque $\mathbf{0}$ si può definire un punto di sella. Nelle nuove coordinate $\mathbf{z} = M\mathbf{y}$, è immediato verificare che gli spazi stabile e instabile sono, rispettivamente, $F^s = \text{span}\{\mathbf{e}^3\}$ e $F^i = \text{span}\{\mathbf{e}^1, \mathbf{e}^2\}$. Dunque, nelle coordinate di partenza, ad esempio, $E^i = MF^i = \text{span}\{\mathbf{m}^1, \mathbf{m}^2\}$.

Presento infine due facili esempi che dovrebbero chiarire l'unico caso finora escluso, quando cioè la matrice A non è invertibile. Anche questo caso è sensibile rispetto a piccole perturbazioni dei coefficienti.

Esempio 3.18. Sia dato il sistema

$$\mathbf{y}' = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \mathbf{y}$$

e proviamo ad analizzare il comportamento delle traiettorie in un intorno U dell'origine. Comunque si scelga U , esistono dati $\mathbf{a} = (a_1, a_2) \in U$ tali che le corrispondenti traiettorie $\mathbf{y}_{\mathbf{a}}$ divergono e dati $\mathbf{a}$ tali che $\mathbf{y}_{\mathbf{a}}(t) \equiv \mathbf{a}$ per ogni t (ovviamente si tratta dei vettori $\mathbf{a}$ tali che $a_1 = 0$). Dunque, $E^i = \text{span}\{\mathbf{e}^1\}$, $E^c = \text{span}\{\mathbf{e}^2\}$. Dunque abbiamo una retta (l'asse y) di punti critici e le traiettorie sono semirette orizzontali che si allontanano da questa. In un certo senso il comportamento è quello di un'equazione scalare.

Esempio 3.19. Sia dato il sistema

$$\mathbf{y}' = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \mathbf{y};$$

qui abbiamo l'autovalore 0 di molteplicità algebrica 2 e geometrica 1. L'autospazio associato a 0 è l'asse delle x , che è luogo di punti critici. Partendo da un dato iniziale (x_0, y_0) , si calcola facilmente che $(x, y)(t) = (x_0 + ty_0, y_0)$; dunque le traiettorie viaggiano su semirette *parallele* all'asse dei punti critici.

3.2 Stabilità dei sistemi non lineari

Torniamo a considerare il problema di Cauchy (PCA): come nel caso lineare, se $\bar{\mathbf{y}} \in \Omega$ è un punto di equilibrio (ossia verifica $\mathbf{f}(\bar{\mathbf{y}}) = 0$), allora la funzione $\mathbf{y}(t) \equiv \bar{\mathbf{y}}$ è una soluzione di

$$\mathbf{y}' = \mathbf{f}(\mathbf{y}) = 0, \quad \mathbf{y}(0) = \bar{\mathbf{y}} \quad (3.9)$$

definita su tutto $[0, +\infty)$. Inoltre, per l'unicità locale, nessun'altra soluzione dell'equazione in (PCA) può assumere in alcun istante il valore $\bar{\mathbf{y}}$. Si noti che l'informazione secondo cui il tempo finale di $\bar{\mathbf{y}}$ è $+\infty$ è ora rilevante, perchè a priori le generiche soluzioni massimali di (PCA) potrebbero esplodere in tempi finiti. Facciamo subito un esempio sul quale torneremo a più riprese anche nel seguito.

Esempio 3.20. Sia $N = 1$ e consideriamo il generico p.d.C. in avanti per l'equazione

$$y' = f(y) = y^3 - y. \quad (3.10)$$

È facile vedere che tale equazione ammette tre punti di equilibrio $y = -1, 0, 1$, corrispondenti agli zeri di f . Supponiamo ad esempio che sia $y_0 \in (0, 1)$. Usando la Prop. 1.22, si dimostra che l'intervallo massimale di esistenza della soluzione corrispondente è $[0, +\infty)$. Per capire il comportamento di y all'infinito è sufficiente allora osservare che $y'(t) < 0$ per ogni t ; dunque, per monotonia, esiste il $\lim_{t \rightarrow +\infty} y(t) =: \ell$. Dal momento che ℓ deve essere un punto di equilibrio, otteniamo subito che $\ell = 0$; la possibilità $\ell = 1$ è infatti esclusa ancora grazie al segno della derivata. Con argomenti analoghi si mostra che, se $y_0 \in (-1, 0)$, allora $\lim_{t \rightarrow +\infty} y(t) = 0$ e se $y_0 \in (1, +\infty)$, allora y esplode a $+\infty$ in tempi finiti.

Dunque, anche nel caso non lineare, certe soluzioni del sistema dinamico (PCA) tendono ad avvicinarsi agli stati di equilibrio. Tuttavia, questo fatto non è sempre vero. Guardando l'esempio precedente, osserviamo che il punto $y = 0$ tende ad "attrarre" le traiettorie che partono vicino ad esso; questo non avviene, invece, per il punto $y = 1$. Infatti, sia per $y_0 > 1$, sia per $y_0 \in (0, 1)$, le traiettorie si vanno allontanando da $y = 1$ man mano che t cresce. Per capire il motivo di questo fatto, vediamo di ricondurre lo studio del sistema (PCA) al caso lineare (PCAL). Il trucco che andiamo a sviluppare si chiama, in effetti, *metodo di linearizzazione* ed è in grado di darci, in molti casi, buone informazioni sul comportamento delle traiettorie. Per introdurlo, abbiamo però bisogno di alcuni preliminari.

Innanzitutto, consideriamo un punto di equilibrio $\bar{\mathbf{y}}$ per il sistema (PCA). Anzi, osserviamo che è ancora possibile supporre $\bar{\mathbf{y}} = \mathbf{0}$; altrimenti, basta effettuare una traslazione. Vogliamo dare delle condizioni sotto le quali $\mathbf{0}$ ha un comportamento riconducibile alla casistica considerata nel caso lineare. L'idea di base è quella di sviluppare $\mathbf{f}$ nell'intorno di $\mathbf{0}$. Grazie alla formula di Taylor, si ha allora

$$\mathbf{f}(\mathbf{y}) = \mathbf{f}(\mathbf{0}) + D\mathbf{f}(\mathbf{0})\mathbf{y} + o(|\mathbf{y}|) = D\mathbf{f}(\mathbf{0})\mathbf{y} + o(|\mathbf{y}|) \quad (3.11)$$

per $\mathbf{y} \rightarrow \mathbf{0}$ (con un piccolo abuso di linguaggio continueremo a usare nel seguito la notazione o piccolo senza più indicare esplicitamente la convergenza $\mathbf{y} \rightarrow \mathbf{0}$), ove $D\mathbf{f}$ denota la matrice Jacobiana di $\mathbf{f}$. Dunque, trascurando l' o piccolo, ci si può aspettare che le soluzioni del sistema si comportino in modo simile a quelle del sistema linearizzato $\mathbf{y}' = A\mathbf{y}$, ove abbiamo posto, qui e nel seguito, $A := D\mathbf{f}(\mathbf{0})$.

Cominciamo ad osservare una proprietà interessante:

Lemma 3.21. *Supponiamo che A sia invertibile. Allora $\mathbf{0}$ è un punto critico isolato per (PCA), ossia esiste un intorno $V \in \mathcal{U}(\mathbf{0})$ tale che $\mathbf{f}(\mathbf{y}) \neq \mathbf{0}$ per ogni $\mathbf{y} \in V \setminus \{\mathbf{0}\}$.*

Prova. Procediamo per assurdo e supponiamo dunque che esista una successione $\{\mathbf{y}_n\} \subset \Omega \setminus \{\mathbf{0}\}$ di punti critici convergente a $\mathbf{0}$. Lo sviluppo di Taylor ci dice allora che

$$\mathbf{0} = \mathbf{f}(\mathbf{y}_n) - \mathbf{f}(\mathbf{0}) = A\mathbf{y}_n + o(|\mathbf{y}_n|).$$

Dividendo tale espressione per $|\mathbf{y}_n|$, otteniamo allora

$$A \frac{\mathbf{y}_n}{|\mathbf{y}_n|} = \frac{o(|\mathbf{y}_n|)}{|\mathbf{y}_n|}. \quad (3.12)$$

Dal momento che $\{\mathbf{y}_n/|\mathbf{y}_n|\}$ è una successione di versori, possiamo estrarne una sotto-successione $\{\mathbf{y}_{n_k}/|\mathbf{y}_{n_k}|\}$ convergente a un versore $\boldsymbol{\xi}$ per $k \rightarrow \infty$. Riscrivendo la (3.12) con n_k al posto di n e passando al limite per $k \rightarrow \infty$, otteniamo subito che $A\boldsymbol{\xi} = \mathbf{0}$; dunque, $\boldsymbol{\xi}$ è un vettore non nullo del nucleo di A , il che è assurdo. ■

Quindi, nel caso A sia invertibile, $\mathbf{0}$ è un punto critico isolato. La situazione è sorprendentemente simile al caso lineare, quando avevamo escluso a priori che la matrice A fosse singolare. Quest'analogia motiva, anche per il caso non lineare, la decisione di escludere dalla trattazione i punti critici non isolati e il Lemma precedente ci viene in aiuto per questo scopo.

Non sarà sorprendente che ulteriori informazioni su $A = D\mathbf{f}(\mathbf{0})$ ci diano ulteriori informazioni sulle traiettorie:

Teorema 3.22. *Supponiamo che tutti gli autovalori di A abbiano parte reale strettamente negativa. Allora il sistema (PCA) decade esponenzialmente nell'intorno di $\mathbf{0}$, ossia si ha che*

1. esiste $U \in \mathcal{U}(\mathbf{0})$, $U \subset \Omega$, tale che, per ogni $\mathbf{a} \in U$, il tempo finale di $\mathbf{y}_{\mathbf{a}}$ è $+\infty$; inoltre, $\mathbf{y}_{\mathbf{a}}(t) \in U$ per ogni $t \in [0, +\infty)$;
2. esistono $b, c > 0$ tali che $\forall \mathbf{a} \in U$ si ha

$$|\mathbf{y}_{\mathbf{a}}(t)| \leq ce^{-bt}|\mathbf{a}| \quad \forall t > 0. \quad (3.13)$$

Prova. Linearizziamo il sistema come in (3.11) e cerchiamo di imitare la dimostrazione del Teorema 3.3. Sia dunque $A = D\mathbf{f}(\mathbf{0})$ e, corrispondentemente, sia V l'intorno fornito dal Lemma 3.21. Scegliamo ora $\mathbf{a} \in V$ e notiamo che la traiettoria $\mathbf{y}_{\mathbf{a}}$ esiste (ed è unica) almeno in piccolo e non si annulla mai. Considerando il prodotto scalare $((\cdot, \cdot))$ e la norma $\|\cdot\|$ forniti dal Lemma 3.2 applicato alla matrice A , ed imitando (3.6), abbiamo

$$\frac{d}{dt}\|\mathbf{y}_{\mathbf{a}}\| = \frac{((\mathbf{f}(\mathbf{y}_{\mathbf{a}}), \mathbf{y}_{\mathbf{a}}))}{\|\mathbf{y}_{\mathbf{a}}\|} = \frac{((A\mathbf{y}_{\mathbf{a}} + o(|\mathbf{y}_{\mathbf{a}}|), \mathbf{y}_{\mathbf{a}}))}{\|\mathbf{y}_{\mathbf{a}}\|}. \quad (3.14)$$

Continuando a ragionare come nel Teorema 3.3, troviamo allora $b > 0$ tale che

$$\frac{((A\mathbf{y}_{\mathbf{a}}, \mathbf{y}_{\mathbf{a}}))}{\|\mathbf{y}_{\mathbf{a}}\|} \leq -2b\|\mathbf{y}_{\mathbf{a}}\|,$$

per ogni valore assunto da $\mathbf{y}_{\mathbf{a}}$. D'altro canto, notiamo che esiste $U \in \mathcal{U}(\mathbf{0})$, $U \subset V$, tale che $\forall \mathbf{z} \in U$, si ha $\|o(|\mathbf{z}|)\| \leq b\|\mathbf{z}\|$. Più precisamente, supponiamo che sia $U = B_{\|\cdot\|}(\mathbf{0}, \delta)$ per qualche $\delta > 0$ (cioè che U sia una “sfera” rispetto alla norma indotta dal nuovo prodotto scalare) e scegliamo ora $\mathbf{a} \in U$. Poichè U è aperto, $\mathbf{y}_{\mathbf{a}}$ resta in U almeno per t sufficientemente piccolo. Per tali t segue allora

$$\frac{d}{dt}\|\mathbf{y}_{\mathbf{a}}(t)\| \leq -b\|\mathbf{y}_{\mathbf{a}}\|,$$

da cui

$$\|\mathbf{y}_{\mathbf{a}}(t)\| \leq \|\mathbf{a}\|e^{-bt}. \quad (3.15)$$

In particolare, grazie anche al Teorema 1.22, $\mathbf{y}_a(t)$ è forzato a restare in U per ogni $t \geq 0$; infatti, la (nuova) norma di $\mathbf{y}_a$ non cresce e U è di forma “sferica”. Di conseguenza, la (3.15) continua a valere $\forall t \geq 0$. La tesi segue tornando eventualmente alla norma euclidea (e qui comparirà, come nel caso lineare, la costante c legata al cambiamento di coordinate). ■

Il Teorema precedente dà buone risposte nel caso che più di ogni altro somiglia alla situazione lineare. Tuttavia, notiamo che ora il decadimento esponenziale del sistema (PCA) si ha soltanto per dati iniziali *sufficientemente piccoli*, mentre per il sistema lineare (PCAL) si aveva il decadimento esponenziale di *ogni* traiettoria.

Questo fatto un po' fastidioso è tuttavia sensato per almeno due motivi: innanzitutto, la dimostrazione sfrutta uno sviluppo di Taylor che in generale approssima bene la funzione $\mathbf{f}$ soltanto *vicino a 0*; inoltre, nel caso nonlineare può benissimo succedere, come del resto accade nell'Esempio 3.20, che vi sia più di un punto di equilibrio $\bar{\mathbf{y}}$. “Vicino” a punti di equilibrio diversi tra loro, il comportamento delle soluzioni sarà ovviamente diverso.

Riprendiamo ancora l'Esempio 3.20 e studiamo più in dettaglio l'equilibrio $\bar{y} = 0$. Essendo $f'(0) = -1$, il Teorema ci assicura che esiste un intorno U di 0 tale che per ogni $y_0 \in U$ si ha che la corrispondente soluzione y di (3.10) verifica

$$\lim_{t \rightarrow +\infty} y(t) = 0,$$

come già avevamo determinato. In più, sappiamo ora che il decadimento è di tipo esponenziale, ed il lettore lo può verificare calcolando esplicitamente la soluzione.

Viceversa, esaminiamo l'equilibrio $\bar{y} = 1$. Allora, $f'(\bar{y}) = 2 > 0$ e viene la curiosità di chiedersi se 1 si comporti in modo analogo a una sorgente esponenziale. Un semplice conto mostra allora che, in questo caso, per $y_0 > 1$ l'andamento della traiettoria è sì sopraesponenziale, ma addirittura questa esplode in tempi finiti, dunque non ha neanche senso guardare cosa succede per $t \rightarrow +\infty$. Se invece $y_0 \in (0, 1)$, allora, già sappiamo che $\lim_{t \rightarrow +\infty} y(t) = 0$. Quindi, in questo caso, la soluzione si allontana dal punto di equilibrio 1, ma resta limitata e dunque non diverge affatto esponenzialmente.

Dunque, già questo semplice esempio monodimensionale suggerisce che, nel caso gli autovalori di A ($= D\mathbf{f}(\mathbf{0})$) abbiano tutti quanti parte reale strettamente positiva, il comportamento per tempi grandi delle soluzioni è diverso da quanto accadeva nel caso lineare e, in generale, il metodo di linearizzazione fornisce ben poche informazioni sull'andamento per tempi grandi. Naturalmente questa è una conseguenza del fatto che, a un certo punto, le traiettorie escono dall'intorno in cui il sistema linearizzato è una “buona approssimazione” di quello non lineare. Questo non poteva succedere nel caso degli autovalori negativi, dove il sistema linearizzato forniva un'approssimazione *globale*. Anzi, ricordando la casistica che può presentarsi nel caso lineare, ci possiamo aspettare che il caso descritto dal Teorema 3.22 sia il più fortunato: ad esempio, basta che un solo autovalore della matrice A abbia parte reale strettamente positiva perché il sistema linearizzato possieda una traiettoria che esce prima o poi dall'intorno di “buona approssimazione”.

Vediamo ora altri due esempi particolarmente interessanti:

Esempio 3.23. Consideriamo l'equazione $y' = -y^3/2$ con la condizione $y(0) = y_0$. Risolvendo esplicitamente, si trova

$$y(t) = \frac{y_0}{(1 + y_0^2 t)^{1/2}},$$

per ogni scelta di $y_0 \in \mathbb{R}$. Dunque, si ha che $\lim_{t \rightarrow +\infty} y(t) = 0$ per ogni y_0 . Noto peraltro che $f'(0) = 0$. Si osservi però che il decadimento delle soluzioni $y(t)$ non è di tipo esponenziale.

Esempio 3.24. Consideriamo l'equazione $y' = y^2$ con la condizione $y(0) = y_0$. Risolvendo esplicitamente, si trova $y(t) = y_0/(1 - y_0 t)$, da cui, se $y_0 > 0$, allora $T = 1/y_0$ e per $t \rightarrow T^-$ si ha che $y(t) \rightarrow +\infty$. Invece, se $y_0 < 0$, è $T = +\infty$ e $\lim_{t \rightarrow +\infty} y(t) = 0^-$. Dunque, da ogni intorno di 0 partono sia traiettorie che convergono a 0 sia traiettorie che esplodono in tempi finiti.

Questi esempi mostrano che la classificazione dei punti di equilibrio effettuata nel caso lineare “sta stretta” al sistema nonlineare (PCA). In effetti, le nozioni di *stabilità* e di *instabilità* di un equilibrio naturali per il caso non lineare sono diverse da quelle viste in precedenza. La definizione fondamentale è la seguente:

Definizione 3.25. Sia $\bar{\mathbf{y}}$ un punto di equilibrio per il sistema (PCA). Allora,

(i) $\bar{\mathbf{y}}$ si dice stabile se $\forall U \in \mathcal{U}(\bar{\mathbf{y}})$, esiste $W \in \mathcal{U}(\bar{\mathbf{y}})$ tale che

$$\mathbf{a} \in W \implies \mathbf{y}_{\mathbf{a}}(t) \in U \quad \forall t > 0; \quad (3.16)$$

(ii) $\bar{\mathbf{y}}$ si dice asintoticamente stabile se è stabile e, inoltre, si può scegliere W in modo tale che $\lim_{t \rightarrow +\infty} \mathbf{y}_{\mathbf{a}}(t) = \bar{\mathbf{y}}$ per ogni $\mathbf{a} \in W$;

(iii) $\bar{\mathbf{y}}$ si dice instabile se non è stabile, ossia se $\exists U \in \mathcal{U}(\bar{\mathbf{y}})$ tale che, per ogni $W \in \mathcal{U}(\bar{\mathbf{y}})$, esiste $\mathbf{y}_0 \in W$ tale che la corrispondente soluzione $\mathbf{y}$ verifica $\mathbf{y}(t) \notin U$ per qualche $t > 0$.

Si noti che l'intorno W di cui è richiesta l'esistenza in (i) verifica necessariamente $W \subset U$. Inoltre, è ovvio che tutti gli intorni che compaiono nella definizione precedente devono essere contenuti nel dominio Ω di $\mathbf{f}$. Osserviamo anche che la condizione (i), grazie alla Prop. 1.22, implica in particolare che il tempo finale di ogni soluzione che parte sufficientemente vicino a $\bar{\mathbf{y}}$ è $+\infty$.

A questo punto suggeriamo al lettore di andare a considerare la classificazione dei sistemi lineari bidimensionali data alla fine del paragrafo precedente e di guardare, caso per caso, quando lo $\mathbf{0}$ è stabile, quando è asintoticamente stabile, e quando è instabile (comunque, questo argomento è trattato in dettaglio alla fine di questo paragrafo). In particolare, si osservi che c'è uno ed un solo caso in cui $\mathbf{0}$ è stabile, ma non asintoticamente stabile. Questo tipo di comportamento non può avvenire per le equazioni non lineari di tipo scalare:

Proposizione 3.26. Sia $N = 1$; dunque, sia I un intervallo aperto e sia $f \in C^1(I; \mathbb{R})$. Sia $\bar{y}$ un punto critico isolato e stabile di (PCA). Allora $\bar{y}$ è asintoticamente stabile.

Prova. Osservo che esiste $\delta > 0$ tale che f non cambia segno né in $(\bar{y}, \bar{y} + \delta]$ né in $[\bar{y} - \delta, \bar{y})$; altrimenti, $\bar{y}$ non sarebbe isolato. Dico allora che deve essere

$$f|_{(\bar{y}, \bar{y} + \delta)} < 0 \quad \text{e} \quad f|_{(\bar{y} - \delta, \bar{y})} > 0. \quad (3.17)$$

Supponiamo infatti che sia falsa una delle relazioni qui sopra, ad esempio la prima. Prendiamo allora $U := (\bar{y} - \delta, \bar{y} + \delta)$. Allora, $\forall W \in \mathcal{U}(\bar{y})$, con $W \subset U$, si ha che, per ogni $a \in W$ con $a > \bar{y}$, $y_a(t)$ esce fuori da U per t sufficientemente grande. Dunque, $\bar{y}$ non potrebbe essere stabile. D'altro canto, da (3.17) segue facilmente l'asintotica stabilità di $\bar{y}$. ■

Qualora si abbia il decadimento esponenziale globale del sistema (ossia valgano le ipotesi del Teorema 3.3 nel caso lineare e del Teorema 3.22 nel caso non lineare), siamo dunque di fronte a un tipo particolarmente fortunato di equilibrio asintoticamente stabile, in cui siamo anche in grado di stimare la velocità di convergenza. L'Esempio 3.23 mostra una situazione differente in cui $\mathbf{0}$ è ancora asintoticamente stabile, anche se la convergenza è “più lenta”. Per quanto riguarda l'Esempio 3.24, lo $\mathbf{0}$ risulta essere un equilibrio instabile: da un certo punto di vista si ha una situazione simile a quella dei punti di sella, poiché da ogni intorno di $\mathbf{0}$ partono sia traiettorie convergenti che traiettorie divergenti.

Per concludere il paragrafo, alla luce della Definizione 3.25, vediamo che cosa può dirci il metodo di linearizzazione discutendo i restanti casi possibili sugli autovalori della matrice A . Come abbiamo osservato in precedenza, le informazioni che potremo ottenere avranno carattere solo locale; infatti, una volta che la soluzione si allontana dal punto di equilibrio, il sistema linearizzato non è più una buona approssimazione del sistema di partenza. C'è però anche un altro fattore di cui tenere conto, che è connesso al discorso sulle “piccole perturbazioni” portato avanti nel paragrafo precedente. Il succo del discorso è il seguente: ci è consentito di scegliere un intorno del punto di equilibrio $\bar{y}$, piccolo quanto vogliamo, in cui studiare il sistema linearizzato.

Ribaltando un po' la prospettiva, possiamo fissare l'attenzione sul sistema linearizzato e vedere il sistema di partenza come una piccola perturbazione di questo, tanto più piccola quanto più restringiamo l'intorno. A questo punto appare chiaro che i casi in cui il sistema linearizzato “funziona” sono tutti quelli in cui i coefficienti della matrice A non sono “sensibili” rispetto a piccole perturbazioni, ossia vale la (i) della Def. 3.9. Questo accade, ad esempio, se tutti gli autovalori sono distinti e hanno parte reale diversa da 0: variando di poco i coefficienti di A restiamo nella stessa situazione. Se invece, ad esempio, A ha un autovalore nullo (o con parte reale nulla), un cambiamento arbitrariamente piccolo dei coefficienti può cambiare sostanzialmente il comportamento di A : può ad esempio far comparire autovalori strettamente positivi o strettamente negativi (ovvero con parte reale strettamente positiva, o strettamente negativa). Analogamente, se ci sono, ad esempio, due autovalori reali coincidenti, una piccola perturbazione dei coefficienti può far comparire autovalori complessi coniugati. In questi ultimi casi, il metodo di linearizzazione fornisce soltanto risposte parziali sul comportamento dell'equilibrio $\bar{y}$.

Vale la pena innanzitutto riportare un risultato generale (vedi anche [9, Teor. 2, p. 269]), la cui dimostrazione (che omettiamo) è tecnicamente più complessa di quella del Teorema 3.22; infatti dobbiamo guardare le traiettorie del sistema nonlineare anche

“lontano” dal punto di equilibrio, dove non abbiamo a disposizione l’aiuto fornito dal sistema linearizzato:

Teorema 3.27. *Sia $\bar{\mathbf{y}}$ un equilibrio del sistema (PCA) e sia $A = D\mathbf{f}(\bar{\mathbf{y}})$. Se A ha un autovalore con parte reale strettamente positiva, allora $\bar{\mathbf{y}}$ è instabile.*

In ogni caso, per chiarire meglio la situazione, ci limitiamo come nel caso lineare a soffermarci sul caso bidimensionale. Notiamo anche che certe proprietà che enunceremo e che paiono piuttosto naturali alla luce della discussione fatta non sono invece immediate da dimostrare in modo rigoroso.

Sistemi bidimensionali (non lineari): Supponiamo $N = 2$ e che $\bar{\mathbf{y}}$ sia un punto critico *isolato* per il sistema (PCA). Poniamo $A := D\mathbf{f}(\bar{\mathbf{y}}) \in \mathcal{M}(\mathbb{R}^2)$ e studiamo il sistema linearizzato. Partendo dal comportamento di questo, vediamo di determinare quali informazioni possiamo ottenere per il sistema nonlineare:

- $\mathbf{0}$ è un nodo (proprio) per il sistema linearizzato. Allora $\bar{\mathbf{y}}$ è un nodo (proprio) per il sistema nonlineare. In particolare, se $\mathbf{0}$ è un nodo stabile, allora il comportamento delle traiettorie di (PCA) vicino a $\bar{\mathbf{y}}$ è simile al comportamento delle traiettorie di (PCAL) vicino a $\mathbf{0}$ per ogni tempo t (e, in particolare, il sistema decade esponenzialmente nell’intorno di $\mathbf{0}$). Se invece $\mathbf{0}$ è un nodo instabile (autovalori positivi distinti), allora le traiettorie di (PCA) si comportano come quelle di (PCAL) solo per tempi piccoli, cioè finché siamo sicuri di essere vicini al punto di equilibrio.
- $\mathbf{0}$ è una sella per il sistema linearizzato. Allora $\bar{\mathbf{y}}$ è una sella per il sistema nonlineare; ossia, vicino a $\bar{\mathbf{y}}$ (e dunque in genere solo per tempi piccoli) le traiettorie si comportano in modo simile a quelle del sistema linearizzato; in particolare $\bar{\mathbf{y}}$ è un equilibrio instabile.
- $\mathbf{0}$ è un nodo improprio per il sistema linearizzato. Supponiamo ad esempio di avere autovalori reali coincidenti strettamente positivi. Allora $\bar{\mathbf{y}}$ è di certo instabile (una perturbazione di A non può modificare il segno della parte reale degli autovalori); tuttavia le traiettorie vicino a $\bar{\mathbf{y}}$ possono comportarsi in vari modi: come nel caso del nodo a una tangente, come nel caso del nodo a stella, come nel caso del nodo proprio, oppure come nel caso del fuoco instabile. Infatti, una perturbazione arbitrariamente piccola di A può dare luogo a una qualunque di queste situazioni.
- $\mathbf{0}$ è un fuoco (stabile o instabile) per il sistema linearizzato. Allora $\bar{\mathbf{y}}$ è un fuoco (rispettivamente stabile o instabile) per il sistema non lineare.
- $\mathbf{0}$ è un centro per il sistema linearizzato. Allora $\bar{\mathbf{y}}$ o è un centro per il sistema nonlineare (si noti che in questo caso è un equilibrio stabile, ma non asintoticamente stabile ai sensi della Def. 3.25), oppure è un fuoco (stabile o instabile) per il sistema non lineare. Infatti, perturbando i coefficienti di A , la parte reale degli autovalori può acquisire un segno.

- ultimo caso, A è una matrice singolare. Allora non si può dire nulla sul sistema nonlineare e, anzi, bisogna innanzitutto accertarsi che $\bar{\mathbf{y}}$ sia un punto critico isolato.

3.3 Stabilità secondo Liapounov

A volte per esaminare la stabilità di un sistema dinamico nonlineare (PCA), l'analisi degli autovalori del corrispondente sistema linearizzato (PCAL) non è il metodo più veloce né il più naturale. Inoltre, come abbiamo visto, ci sono casi di indecisione (ad esempio quando la matrice A dei coefficienti ha autovalori nulli o con parte reale nulla). Per costruire una procedura alternativa, riguardiamo la dimostrazione del Teorema 3.22. L'idea che si è usata è piuttosto semplice: si è provato che *una* norma della soluzione decresce (esponenzialmente) lungo le traiettorie. Tuttavia, questo non si riusciva a fare direttamente per la norma euclidea; si andavano a prendere un diverso prodotto scalare e la nuova norma associata a questo. In questo paragrafo generalizziamo ulteriormente un tale approccio; andremo infatti a determinare un'ampia famiglia di funzioni la cui decrescenza lungo le traiettorie continui a garantire la stabilità di un equilibrio.

Uno strumento che useremo più volte nel seguito è il seguente corollario del Teorema 1.28 (e qui si capisce perché abbiamo voluto dimostrare quel teorema sotto le ipotesi d'esistenza in piccolo):

Corollario 3.28. *Sia data $\mathbf{f} \in C^1(\Omega; \mathbb{R}^N)$ e si consideri la famiglia di problemi di Cauchy*

$$\mathbf{y}'_n = \mathbf{f}(\mathbf{y}_n), \quad \mathbf{y}_n(0) = \mathbf{y}_0^n, \quad (3.18)$$

ove $\mathbf{y}_0^n \rightarrow \mathbf{y}_0 \in \Omega$ per $n \rightarrow \infty$. Sia $\mathbf{y}$ la soluzione massimale del problema di Cauchy limite

$$\mathbf{y}' = \mathbf{f}(\mathbf{y}), \quad \mathbf{y}(0) = \mathbf{y}_0. \quad (3.19)$$

Allora, per ogni $t \in \text{dom } \mathbf{y}$, si ha che $\mathbf{y}_n$ è definita in t almeno per n sufficientemente grande; inoltre, si ha che $\mathbf{y}_n(t)$ tende a $\mathbf{y}(t)$.

Questo corollario ammette un'interpretazione interessante, per la quale vale la pena di introdurre una nuova notazione. Al variare di $t \geq 0$ (ma a volte ammetteremo anche tempi negativi), associamo al sistema dinamico (PCA) l'operatore-soluzione (o *semigrupp*)

$$S(t) : \Omega \rightarrow \Omega, \quad S(t) : \mathbf{y}_0 \mapsto \mathbf{y}(t). \quad (3.20)$$

Si osservi che l'operatore $S(t)$ ha senso soltanto per i tempi t tali che $t \in \text{dom}(\mathbf{y})$, dove $\text{dom}(\mathbf{y})$ dipende a priori dal dato $\mathbf{y}_0$. Con una tale restrizione sui tempi ammissibili, il corollario dice essenzialmente che l'operatore $S(t)$ è continuo per ogni $t \in \text{dom } \mathbf{y}$.

Notiamo peraltro che, nei casi interessanti per l'analisi della stabilità, si ha che in generale la soluzione è definita su tutto $\mathbb{R}^+$: in tal caso, $S(t)$ risulta definito e continuo per ogni $t \geq 0$. Altre proprietà naturali e importanti di $S(t)$ sono le seguenti:

- (S1) $S(0) = \text{Id}$, applicazione identica;
- (S2) $S(t+s) = S(t) \circ S(s)$ per ogni $t, s \geq 0$.

Si noti anche che il concetto di semigruppato si può adattare ed estendere a situazioni molto più generali di quelle affrontate in questi appunti. In tale ottica, una famiglia $\{S(t)\}_{t \geq 0}$ di mappe definite su un insieme X (che supponiamo abbia una struttura topologica ragionevole) e a valori in X stesso si dice semigruppato se soddisfa le proprietà (S1) e (S2). Si parla di solito di semigruppato *continuo* quando valgono la (3.20) e la seguente

$$\forall x \in X, \quad t \mapsto S(t)x \quad \text{è una funzione continua.} \quad (3.21)$$

Nel nostro caso, la (3.21) corrisponde a dire che per ogni scelta del dato iniziale $\mathbf{a} \in \Omega$, la corrispondente soluzione

$$t \mapsto S(t)\mathbf{a} = \mathbf{y}_{\mathbf{a}}(t)$$

è una funzione continua da $[0, +\infty)$ in Ω (e questo è vero; anzi, sappiamo che è addirittura di classe C^1). Possiamo allora dire che la mappa $S(t)$ associata al sistema dinamico (PCA) è un semigruppato continuo purché ogni soluzione sia definita su $[0, +\infty)$. Altrimenti, le proprietà (S1), (S2), (3.20), (3.21) valgono comunque con le opportune restrizioni su t .

Possiamo ora passare al risultato principale della cosiddetta teoria di Liapounov.

Teorema 3.29. *Sia $\mathbf{f} \in C^1(\Omega; \mathbb{R}^N)$ e sia $\bar{\mathbf{y}} \in \Omega$ un punto di equilibrio isolato per il sistema (PCA), ove $\mathbf{f} \in C^1(\Omega; \mathbb{R}^N)$. Supponiamo che esistano $\Lambda \in \mathcal{U}(\bar{\mathbf{y}})$, $\Lambda \subset \Omega$ ed una funzione $V : \Lambda \rightarrow \mathbb{R}$ tale che:*

$$V \in C^0(\Lambda), \quad V \text{ differenziabile in } \Lambda \setminus \{\bar{\mathbf{y}}\}, \quad (3.22)$$

$$V(\bar{\mathbf{y}}) = 0, \quad V > 0 \quad \text{in } \Lambda \setminus \{\bar{\mathbf{y}}\}, \quad (3.23)$$

$$\dot{V}(\mathbf{y}) := \nabla V(\mathbf{y}) \cdot \mathbf{f}(\mathbf{y}) \leq 0 \quad \forall \mathbf{y} \in \Lambda \setminus \{\bar{\mathbf{y}}\}. \quad (3.24)$$

Allora $\bar{\mathbf{y}}$ è un punto di equilibrio stabile per (PCA) e V si dice funzione di Liapounov associata al sistema (PCA) e all'equilibrio $\bar{\mathbf{y}}$. Inoltre, se oltre a (3.22–3.24) vale la

$$\dot{V}(\mathbf{y}) < 0 \quad \forall \mathbf{y} \in \Lambda \setminus \{\bar{\mathbf{y}}\}, \quad (3.25)$$

allora $\bar{\mathbf{y}}$ è asintoticamente stabile e V si dice funzione di Liapounov stretta.

Prova. Verifichiamo la proprietà (i) della Definizione 3.25. Sia dunque U come in (i) e, anzi, eventualmente rimpiccioliamo U in modo tale che sia

$$U = B(\bar{\mathbf{y}}, \delta) \quad \text{per qualche } \delta > 0; \quad \bar{B}(\bar{\mathbf{y}}, \delta) \subset \Lambda. \quad (3.26)$$

In questo modo, siamo certi che V è definita su tutto il nostro U , e addirittura sulla sua chiusura. Sia inoltre $\mathbf{a} \in U \setminus \{\bar{\mathbf{y}}\}$ e sia $\mathbf{y}_{\mathbf{a}}$ la traiettoria del sistema che parte da $\mathbf{a}$. Osserviamo innanzitutto che per t sufficientemente piccolo deve essere $\mathbf{y}_{\mathbf{a}}(t) \in U \setminus \{\bar{\mathbf{y}}\}$ e dunque, grazie a (3.24),

$$0 \geq \dot{V}(\mathbf{y}_{\mathbf{a}}(t)) = \nabla V(\mathbf{y}_{\mathbf{a}}(t)) \cdot \mathbf{y}'_{\mathbf{a}}(t) = \frac{d}{dt} V(\mathbf{y}_{\mathbf{a}}(t)); \quad (3.27)$$

dunque, effettivamente $\dot{V}$ può essere interpretata come la derivata di V lungo una generica traiettoria. Definiamo ora α come il minimo di V sul compatto ∂U (frontiera

della bolla $B(\bar{\mathbf{y}}, \delta)$). Grazie a (3.23), si ha $\alpha > 0$ e possiamo definire $W := \{\mathbf{y} \in U : V(\mathbf{y}) < \alpha\}$. Poiché V è continua, W è aperto e contiene $\bar{\mathbf{y}}$ grazie alla prima delle (3.23). Inoltre, per definizione, $W \subset U$. Restringiamoci allora a prendere $\mathbf{a} \in W$. Grazie a (3.27), è evidente che $\mathbf{y}_{\mathbf{a}}$ non può mai uscire da U (anzi, neanche da W). Questo mostra che $\bar{\mathbf{y}}$ è stabile.

Supponiamo ora che V sia una funzione di Liapounov stretta e sia $\mathbf{a} \in W$. Dobbiamo mostrare che $\exists \ell = \lim_{t \rightarrow +\infty} S(t)\mathbf{a}$ e che $\ell = \bar{\mathbf{y}}$. Per fare questo, procediamo per assurdo. Nel seguito, indicheremo con la notazione $\{t_n \nearrow\}$ una successione $\{t_n\} \subset [0, +\infty)$ che sia strettamente crescente e divergente. Poiché $\mathbf{y}_{\mathbf{a}}([0, +\infty)) \subset \bar{B}(\bar{\mathbf{y}}, \delta)$, che è chiuso e limitato, se la tesi non vale esistono una successione $\{t_n \nearrow\}$ ed un punto $\mathbf{z} \neq \bar{\mathbf{y}}$, $\mathbf{z} \in \bar{B}(\bar{\mathbf{y}}, \delta)$ tali che

$$\lim_{n \rightarrow \infty} S(t_n)\mathbf{a} = \mathbf{z}.$$

Anzi, grazie alla (3.24), si ha che $\mathbf{z} \in W$. Consideriamo allora la soluzione $\mathbf{y}_{\mathbf{z}} = S(\cdot)\mathbf{z}$ di (PCA) con dato iniziale $\mathbf{z}$. Poiché $\mathbf{z} \in W$, certamente $\mathbf{y}_{\mathbf{z}}$ è definita per ogni $t \geq 0$. Applicando il Cor. 3.28 e ricordando la proprietà (S2), $\forall T > 0$ otteniamo

$$S(T + t_n)\mathbf{a} = S(T)S(t_n)\mathbf{a} \rightarrow S(T)\mathbf{z}. \quad (3.28)$$

D'altro canto, poichè V è *strettamente* decrescente lungo le traiettorie, valgono le relazioni

$$V(\mathbf{z}) < V(S(t)\mathbf{a}), \quad V(S(t)\mathbf{z}) < V(\mathbf{z}), \quad \forall t > 0 \quad (3.29)$$

(la prima di queste è vera poiché $\forall t > 0$ esiste $t_n > t$; dunque $V(\mathbf{z}) < V(S(t_n)\mathbf{a}) < V(S(t)\mathbf{a})$). Ora, la successione $n \mapsto V(S(t_n + T)\mathbf{a})$ è strettamente decrescente rispetto a n e dunque ammette limite. Inoltre,

$$\begin{aligned} V(\mathbf{z}) &\leq \lim_{n \rightarrow \infty} V(S(t_n + T)\mathbf{a}) \quad (\text{per la prima delle (3.29)}) \\ &= V(S(T)\mathbf{z}) \quad (\text{per la (3.28) e la continuità di } V) \\ &< V(\mathbf{z}) \quad (\text{per la seconda delle (3.29)}), \end{aligned} \quad (3.30)$$

il che è assurdo. ■

Osservazione 3.30. Si noti che la funzione $\dot{V}$ introdotta in (3.24) dipende a priori solo da V e da $\mathbf{f}$ e non dal tempo. Tuttavia, nella dimostrazione del Teorema si è visto che, ogniquale volta calcoliamo $\dot{V}$ in un punto del tipo $\mathbf{y}(t)$, ove $\mathbf{y}$ è una traiettoria, $\dot{V}(\mathbf{y}(t))$ coincide proprio con la derivata in tempo di $V \circ \mathbf{y}$. Pertanto, la funzione $\dot{V}$ si chiama a volte derivata di V *lungo le traiettorie*.

Vediamo di commentare in breve il risultato precedente. La grande generalità con cui possiamo scegliere la funzione V è compensata dalle minori informazioni che il metodo fornisce: non possiamo dire se $|\mathbf{y}_{\mathbf{a}} - \bar{\mathbf{y}}|$ decresce esponenzialmente e, come si è visto, già la dimostrazione della stabilità asintotica (che è una proprietà ben più debole) è piuttosto laboriosa.

Per presentare qualche esempio di applicazione del teorema, osserviamo che la definizione di funzione di Liapounov è naturale non soltanto dal punto di vista matematico: in molti casi “pratici”, la V è una qualche grandezza caratteristica del sistema,

ad esempio un'energia, la cui decrescenza sia dettata da un principio fisico (questo punto sarà ulteriormente precisato nel seguito). Ovviamente, questa considerazione può dare un aiuto nel cercare una funzione di Liapounov per un determinato sistema, dal momento che il Teorema precedente non fornisce alcun procedimento costruttivo.

Esempio 3.31. Consideriamo un punto materiale di massa unitaria che si muove sotto l'azione di un campo di forze conservativo $\mathbf{f} = -\nabla\Phi$ definito (e sufficientemente regolare) in un aperto $\Omega \subset \mathbb{R}^3$. Denotiamo con $\mathbf{x}$ la posizione e con $\mathbf{v}$ la velocità del punto. Per la legge di Newton, otteniamo il sistema

$$\mathbf{x}' = \mathbf{v}, \quad \mathbf{v}' = \mathbf{f}(\mathbf{x}). \quad (3.31)$$

Forse vale la pena osservare che (3.31) “vive” in $\mathbb{R}^6$; in effetti, l'incognita è $\mathbf{y} = (\mathbf{x}, \mathbf{v})$. Osserviamo che $(\bar{\mathbf{x}}, \bar{\mathbf{v}})$ è un equilibrio per (3.31) se e solo se $\bar{\mathbf{v}} = \mathbf{0}$ e $\nabla\Phi(\bar{\mathbf{x}}) = \mathbf{0}$. Vediamo sotto quali condizioni l'energia totale del sistema è un funzionale di Liapounov; più precisamente, affinché sia soddisfatta (3.23), dobbiamo porre

$$V(\mathbf{x}, \mathbf{v}) := \frac{1}{2}|\mathbf{v}|^2 + \Phi(\mathbf{x}) - \Phi(\bar{\mathbf{x}}).$$

Allora, si ha subito

$$\dot{V}(\mathbf{x}, \mathbf{v}) = \nabla_{\mathbf{x}}V \cdot \mathbf{x}' + \nabla_{\mathbf{v}}V \cdot \mathbf{v}' = -\mathbf{v}' \cdot \mathbf{v} + \mathbf{v} \cdot \mathbf{v}' \equiv 0,$$

come d'altronde richiede il principio di conservazione dell'energia. Dunque, valgono (3.22) e (3.24); affinché sia verificata (3.23), dobbiamo supporre che $\Phi(\mathbf{x}) > \Phi(\bar{\mathbf{x}})$ per ogni $\mathbf{x} \neq \bar{\mathbf{x}}$ in un intorno di $\bar{\mathbf{x}}$; ossia, $\bar{\mathbf{x}}$ deve essere un minimo locale stretto per il potenziale Φ .

L'esempio introdotto qui sopra è particolarmente interessante anche per un motivo matematico. Notiamo infatti che la funzione V verifica non solo la (3.24), ma addirittura l'uguaglianza $\dot{V} = 0$ in ogni punto. Una funzione che ha questa proprietà in tutto Ω (dunque indipendentemente dalla presenza o meno di punti di equilibrio) viene normalmente indicata nel seguente modo.

Definizione 3.32. Si consideri il sistema (PCA) associato ad una funzione $\mathbf{f} \in C^1(\Omega; \mathbb{R}^N)$. Una funzione $E \in C^1(\Omega)$ si dice integrale primo del sistema (PCA) qualora si abbia

$$\dot{E}(\mathbf{y}) = \nabla E(\mathbf{y}) \cdot \mathbf{f}(\mathbf{y}) = 0 \quad \forall \mathbf{y} \in \Omega. \quad (3.32)$$

A cosa servono gli integrali primi? L'osservazione fondamentale è che, grazie alla (3.32), l'immagine di ogni traiettoria di un sistema (PCA) che ammette un integrale primo E è contenuta in un insieme di livello di E . Questa è un'informazione molto forte. Ad esempio, in dimensione 2, gli insiemi di livello di E sono spesso (sostegni di) curve e questo fatto di solito permette di identificare completamente le traiettorie. In dimensione superiore le informazioni sono un po' minori, ma sempre rilevanti. Ad esempio, se un integrale primo ha tutti gli insiemi di livello limitati, ogni traiettoria del sistema è definita su tutto $\mathbb{R}$ (dimostrare per esercizio!). Nell'Esempio 3.31 (e ogniquale volta un punto di equilibrio ammette una funzione di Liapounov che è anche un

integrale primo del sistema), vediamo allora direttamente che i punti di equilibrio del sistema (3.31) sono tutti stabili, ma non asintoticamente stabili. In generale notiamo comunque che, specialmente in dimensione maggiore di 2, è tutt'altro che scontato che un sistema dinamico ammetta un integrale primo.

Questo discorso ci aiuterà anche ad enunciare un nuovo criterio per la stabilità asintotica, utile anche nelle applicazioni e negli esercizi, e che vale in certi casi in cui un punto di equilibrio ammette una funzione di Liapounov, ma questa non è una funzione di Liapounov stretta (ossia può valere l'uguale in (3.24)). Per introdurre questo Teorema, ci servono tuttavia alcune nuove definizioni, che avranno peraltro importanza anche nel prossimo paragrafo:

Definizione 3.33. Sia $\mathbf{a} \in \Omega$. Si chiama ω -limite di $\mathbf{a}$ l'insieme

$$\omega\text{-lim } \mathbf{a} := \{ \mathbf{z} \in \Omega : \exists \{t_n \nearrow\} : S(t_n)\mathbf{a} \rightarrow \mathbf{z} \}. \quad (3.33)$$

Osservazione 3.34. Se la traiettoria che origina da $\mathbf{a}$ esplode in tempi finiti, in genere si dice che l' ω -limite di $\mathbf{a}$ è vuoto; naturalmente si ha che l' ω -limite è vuoto anche quando la traiettoria ha tempo finale $+\infty$ ma è divergente.

Osservazione 3.35. Analogamente, si dà la definizione di α -limite sostituendo $t_n \nearrow +\infty$ con $t_n \searrow -\infty$. È chiaro che questa terminologia si riferisce al fatto che α e ω sono, rispettivamente, la prima e l'ultima lettera dell'alfabeto greco.

Si noti peraltro che nella teoria che introdurremo supporremo spesso che il tempo finale di una soluzione sia $+\infty$, mentre non faremo nessun'ipotesi sul tempo iniziale; dunque non sempre potremo definire l' α -limite.

Definizione 3.36. Diciamo che un insieme $P \subset \Omega$ è positivamente invariante per (PCA) se $\mathbf{a} \in P$ implica che $S(t)\mathbf{a}$ è definito e appartiene a P per ogni $t \geq 0$.

Più succintamente avremmo potuto dire " $S(t)P \subset P$ per ogni $t > 0$ ". Notiamo che è parte della definizione il fatto che ogni traiettoria che origina da $\mathbf{a} \in P$ abbia tempo finale $+\infty$. Analogamente,

Definizione 3.37. Diciamo che un insieme $P \subset \Omega$ è invariante ovvero completamente invariante per (PCA) se $S(t)P = P$ per ogni $t \geq 0$.

Osservazione 3.38. Si noti che, se P è completamente invariante, allora ogni traiettoria $\mathbf{y}_{\mathbf{a}}$ con $\mathbf{a} \in P$ è definita su tutto $\mathbb{R}$. Inoltre, P è completamente invariante se e solo se è sia positivamente che negativamente invariante. Questa seconda condizione significa che $S(-t)P \subset P$ per ogni $t > 0$ (ossia, per ogni $\mathbf{a} \in P$, $S(-t)\mathbf{a} := \mathbf{y}_{\mathbf{a}}(-t)$ è definita e sta in $P \forall t > 0$). Infine, notiamo che, se P è completamente invariante, la mappa $S(t) : P \rightarrow P$ è biunivoca per ogni $t \in \mathbb{R}$. Invitiamo il lettore a dimostrare in dettaglio tutte queste proprietà. Gli strumenti fondamentali che devono essere usati sono il Teorema di esistenza e unicità in piccolo e la proprietà (S2).

Definizione 3.39. Chiamiamo orbita del sistema (PCA) l'immagine di una traiettoria $\mathbf{y}$, ossia l'insieme $\mathbf{y}(\text{dom}(\mathbf{y}))$. Diciamo che una traiettoria $\mathbf{y}$ individua un'orbita intera (o, a volte, con abuso di linguaggio, che è un'orbita intera), se $\mathbf{y}$ è una soluzione definita per ogni $t \in \mathbb{R}$.

Esercizio 3.40. Quando accade che due traiettorie individuano la stessa orbita?

Si noti che una traiettoria individua un'orbita intera se e solo se tale orbita è un insieme completamente invariante. Abbiamo dunque costruito una famiglia particolarmente naturale di insiemi completamente invarianti.

Veniamo ora al nuovo criterio per la stabilità asintotica:

Teorema 3.41. *Sia $\bar{\mathbf{y}}$ un equilibrio stabile per il sistema (PCA). Esista inoltre un insieme $P \in \mathcal{U}(\bar{\mathbf{y}})$ chiuso, limitato, positivamente invariante e contenuto in Ω . Esista inoltre una funzione di Liapounov V associata all'equilibrio $\bar{\mathbf{y}}$ e definita in P . Infine, supponiamo che (a parte $\mathbf{y} \equiv \bar{\mathbf{y}}$) P non contenga alcun'orbita intera del sistema su cui V è costante. Allora $\bar{\mathbf{y}}$ è asintoticamente stabile.*

Osservazione 3.42. Prima di fare la dimostrazione è opportuno spendere due parole sull'ultima ipotesi. Se, in particolare, siamo capaci di dimostrare che P non contiene alcun'orbita intera, senza ulteriori specificazioni, tanto meglio. In effetti, come è detto nell'enunciato (e si vedrà nella dimostrazione), è sufficiente la condizione (più debole!) che P non contenga orbite intere che vivono sugli insiemi di livello di V .

Prova. Mostro che $\forall \mathbf{a} \in P$, $S(t)\mathbf{a}$ tende a $\bar{\mathbf{y}}$, cosicchè (ii) di Def. 3.25 è soddisfatta con $W = P$. Procediamo per assurdo e supponiamo che esista $\mathbf{a} \in P$ tale che $S(t)\mathbf{a}$ non tende a $\bar{\mathbf{y}}$. Questo vale a dire che $\omega\text{-lim } \mathbf{a}$ contiene punti diversi da $\bar{\mathbf{y}}$. Notiamo peraltro che, dati due qualunque punti $\mathbf{z}_1, \mathbf{z}_2 \in \omega\text{-lim } \mathbf{a}$, si ha che

$$V(\mathbf{z}_1) = V(\mathbf{z}_2) = \inf_{t \geq 0} V(S(t)\mathbf{a}) =: \lambda; \quad (3.34)$$

in particolare, $\omega\text{-lim } \mathbf{y}_{\mathbf{a}}$ è contenuto in un insieme di livello di V (il lettore dimostri in dettaglio questa proprietà).

Ora, la negazione della tesi assicura che esiste almeno un punto $\mathbf{z} \neq \bar{\mathbf{y}}$ appartenente a $\omega\text{-lim } \mathbf{a}$. Poichè P è chiuso e positivamente invariante, otteniamo subito che $\mathbf{z} \in P$; inoltre, grazie alla (3.34), il valore $\lambda = V(\mathbf{z})$ è indipendente dalla scelta di $\mathbf{z}$ in $\omega\text{-lim } \mathbf{a}$. Sia allora $\{t_n \nearrow\}$ una successione tale che $S(t_n)\mathbf{a}_n \rightarrow \mathbf{z}$.

Considero ora la traiettoria $\mathbf{y}_{\mathbf{z}}$, soluzione massimale di (PCA) con dato iniziale $\mathbf{z}$ e chiamo $T^- < 0$ e $T^+ > 0$ i suoi tempi iniziale e, rispettivamente, finale. Voglio mostrare che $T^- = -\infty$ e $T^+ = +\infty$. Che sia $T^+ = +\infty$ è chiaro; infatti $\mathbf{z}$ sta nell'insieme positivamente invariante P . D'altro canto, grazie ad (S2) ed al Coroll. 3.28, per ogni $t \in (T^-, +\infty)$ si ha

$$S(t)\mathbf{z} = \lim_{n \rightarrow \infty} S(t)S(t_n)\mathbf{a} = \lim_{n \rightarrow \infty} S(t_n + t)\mathbf{a} \in \omega\text{-lim } \mathbf{a} \subset P. \quad (3.35)$$

Segue allora che l'immagine di $(T^-, +\infty)$ tramite $\mathbf{y}_{\mathbf{z}}$ è contenuta nel sottoinsieme compatto P di Ω . Grazie al Teorema 1.22, si ha allora che $T^- = -\infty$; dunque, il punto $\mathbf{z}$ individua un'orbita intera.

A questo punto, si può ripetere il conto (3.35) per ogni $t \in \mathbb{R}$. Questo ci dice che $S(t)\mathbf{z} \in \omega\text{-lim } \mathbf{a} \forall t \in \mathbb{R}$; dunque, grazie a (3.34), abbiamo in particolare che $V(\mathbf{y}_{\mathbf{z}}(t)) = \lambda$ per ogni $t \in \mathbb{R}$ e questo dà l'assurdo desiderato. ■

Vediamo ora di chiarire quale può essere nei casi concreti l'utilità del precedente Teorema, che ha un enunciato decisamente astratto. Si supponga di stare studiando un

sistema del tipo (PCA) e di avere costruito una funzione di Liapounov V relativa a un equilibrio $\bar{\mathbf{y}}$. Per dimostrare che $\bar{\mathbf{y}}$ è asintoticamente stabile, la prima cosa da fare è vedere se vale la condizione (3.25). Sfortunatamente, però, può capitare che $\dot{V}$ abbia degli zeri e non sia una funzione di Liapounov stretta. Tuttavia, supponendo per semplicità $\Omega = \mathbb{R}^N$, in molti casi concreti la funzione V verifica una relazione di coercività del tipo

$$\lim_{|\mathbf{y}| \rightarrow +\infty} V(\mathbf{y}) = +\infty. \quad (3.36)$$

In tal caso, per ogni $\zeta > 0$, gli insiemi

$$P_\zeta := V^{-1}([0, \zeta])$$

risultano essere intorno di $\bar{\mathbf{y}}$ chiusi, limitati (grazie a (3.36)) e positivamente invarianti (grazie a (3.24)). Se V è vista come un'energia, i V_ζ corrispondono alle configurazioni del sistema con energia minore o uguale a ζ . Allora, grazie al Teorema, se $\bar{\mathbf{y}}$ non è asintoticamente stabile, comunque piccolo noi scegliamo ζ deve per forza esistere un'orbita intera $\mathbf{y}$ contenuta in V_ζ ma che non tende a $\bar{\mathbf{y}}$. In certi casi, uno studio qualitativo delle traiettorie permette di escludere questa possibilità e dunque di concludere a favore della stabilità asintotica del punto critico.

4 Sistemi dinamici, orbite, attrattori

In questo capitolo, presenteremo alcune nozioni più avanzate della teoria dei sistemi dinamici, col fine principale di introdurre l'oggetto chiamato "attrattore globale". Rispetto a quanto visto fino ad ora, smetteremo di pensare direttamente in termini di equazioni, ma adotteremo un approccio più generale ed astratto, che consentirà di applicare la teoria anche a situazioni diverse. L'impostazione "concreta" adottata fino a questo punto dovrà comunque essere tenuta ben presente, in quanto servirà a fornire la maggior parte delle applicazioni pratiche e a chiarire certi concetti astratti altrimenti di difficile interpretazione.

4.1 Un approccio astratto

In questo paragrafo riprenderemo in modo più organico e generale alcuni concetti visti qua e là nelle pagine precedenti, adattandoli al nuovo contesto che vogliamo introdurre. Il primo passo consiste nella seguente definizione di semigruppato, che riformula in termini astratti molte proprietà che abbiamo dimostrato per le soluzioni di un Problema di Cauchy in avanti autonomo.

Definizione 4.1. *Sia $\mathcal{X}$ uno spazio metrico. Si dice semigruppato continuo su $\mathcal{X}$ una famiglia $\{S(t), t \in [0, +\infty)\}$ di applicazioni di $\mathcal{X}$ in sè tali che:*

$$S(0) = \text{Id}, \quad (S1)$$

$$S(t) \text{ è continua } \quad \forall t \geq 0, \quad (S2)$$

$$S(t) \circ S(s) = S(t+s) \quad \forall t, s \geq 0, \quad (S3)$$

$$t \mapsto S(t)x \text{ è continua } \quad \forall x \in \mathcal{X}. \quad (S4)$$

Indicheremo nel seguito un semigruppoo continuo con S o $S(t)$; inoltre, per brevità, quando useremo il termine “semigruppoo” intenderemo sempre implicitamente che questo sia continuo. L’insieme $\mathcal{X}$ si dice spazio delle fasi del semigruppoo.

Tornando all’esempio del Problema di Cauchy, la mappa $S(t)$ associa al generico dato $\mathbf{a} \in \Omega =: \mathcal{X}$ la soluzione $\mathbf{y}_{\mathbf{a}}$ valutata al tempo t . In quest’ottica, la (S1) è legata all’acquisizione del dato iniziale; la (S2) segue dal Teorema di dipendenza continua; la (S3) è di fatto la prolungabilità delle soluzioni; infine la (S4) è legata alla loro regolarità (su questo punto torneremo fra un attimo).

Osservazione 4.2. Per il momento non è necessario chiedere che $\mathcal{X}$ sia completo. Questo fatto ha il vantaggio che possiamo considerare il caso in cui $\mathcal{X} = \Omega$ aperto di $\mathbb{R}^N$ e includere nell’approccio astratto le equazioni autonome su Ω . Vedremo tuttavia nel seguito della teoria che una condizione di completezza (o qualcosa di equivalente) sarà necessaria; dunque il caso $\mathcal{X} = \Omega$ richiederà qualche aggiustamento.

Rispetto a quanto detto prima, stiamo ribaltando il punto di vista. Infatti, ora il punto di partenza è il concetto “astratto” di semigruppoo, e le mappe $S(t)$ potrebbero anche non essere legate ad equazioni differenziali. L’interpretazione “dinamica” viene comunque parzialmente mantenuta. Anzi, useremo nel seguito l’espressione “sistema dinamico” come sinonimo di “semigruppoo”.

Definizione 4.3. Sia S un semigruppoo sullo spazio metrico $\mathcal{X}$. Si dice traiettoria associata a S ogni funzione continua $u : I \rightarrow X$, ove I è un intervallo di $\mathbb{R}$ contenente $[0, +\infty)$, e tale che

$$S(t)u(\tau) = u(t + \tau) \quad \forall \tau \in I, \quad t \geq 0. \quad (4.1)$$

Una traiettoria u si dice massimale se non ammette prolungamenti propri (ossia se non esiste alcuna traiettoria $v : J \rightarrow X$ tale che $I \subset J$, $I \neq J$ e $v|_I \equiv u$); u si dice completa se $I = \mathbb{R}$.

È evidente che ogni traiettoria completa è massimale. Si noti che, analogamente, avremmo potuto dapprima definire che cosa è, in astratto, una traiettoria e poi dare la definizione di semigruppoo. L’approccio usato sembra comunque più naturale; ad esempio, formulare la proprietà (S2) di continuità in termini di traiettorie sarebbe stato abbastanza complicato.

Oltre al fatto di essere in uno spazio metrico qualunque piuttosto che in un aperto di $\mathbb{R}^N$, ci sono almeno vari altri aspetti della teoria che sono più generali rispetto a quanto visto per le equazioni differenziali. Innanzitutto, le traiettorie $u(\cdot)$, che finora erano sempre mappe C^1 , se non C^2 , ora possono essere anche solo C^0 (del resto, non è in generale possibile parlare di derivata di una funzione a valori in uno spazio metrico). Un altro aspetto rilevante, ma più riposto, è il seguente: dato $x \in \mathcal{X}$, $\forall t \geq 0$ è sempre definito $S(t)x$ e, se pensiamo all’esempio delle equazioni differenziali, questo corrisponde alla proprietà di esistenza e unicità (in grande) delle soluzioni per il problema *in avanti*. Tuttavia, potrebbe ora capitare che la mappe $S(t)$ non siano iniettive. Ossia, potrebbe accadere che $S(t)x_1 = S(t)x_2$ per qualche $x_1, x_2 \in \mathcal{X}$ e qualche $t > 0$. Ovviamente una situazione di questo tipo non avviene nel caso delle equazioni differenziali ordinarie, almeno quando ci possiamo appoggiare al teorema

di esistenza e unicità in piccolo, ma capita frequentemente in contesti più generali. Diamo dunque una

Definizione 4.4. *Dico che un semigruppato S verifica la proprietà di unicità all'indietro se $S(t)$ è iniettiva per ogni $t \geq 0$.*

Esempio 4.5. Si consideri il problema di Cauchy

$$y' = -(\sqrt{|y|} + y^2) \operatorname{sign}(y), \quad (4.2)$$

ove sign è la funzione segno. Si vede facilmente che (4.2) genera un semigruppato continuo (in particolare, per il problema *in avanti* si ha unicità della soluzione), per il quale non vale la proprietà di unicità all'indietro. Si noti anche che le soluzioni di (4.2) sono definite $\forall t \geq 0$, ma esplodono in tempi finiti negativi. Questo non è in contrasto con la Def. 4.1.

Come nel caso concreto, possiamo parlare di *orbite* di un sistema dinamico:

Definizione 4.6. *Sia S un semigruppato. Si chiama orbita di S l'immagine di una traiettoria massimale. Un'orbita si dice “completa” o “intera” se è immagine di una traiettoria completa.*

Osservazione 4.7. La definizione data non è universalmente accettata. Su alcuni testi, un'orbita è definita come insiemi dei valori in $\mathcal{X}$ assunti da una traiettoria al variare di $t \in [0, +\infty)$ (vengono cioè esclusi eventuali “tempi di vita” negativi). Noi chiameremo tali insiemi “orbite positive”.

Occorre naturalmente fare distinzione tra le orbite, che sono sottoinsiemi di $\mathcal{X}$, e le traiettorie, che sono funzioni del tempo a valori in $\mathcal{X}$. Si noti (cfr. Esempio 4.5) che non sempre due orbite sono insiemi disgiunti. Questo è invece vero se vale la proprietà di unicità all'indietro; in tale caso, la famiglia delle orbite di S costituisce una *partizione* dello spazio delle fasi $\mathcal{X}$.

Osservazione 4.8. Rispetto a quanto visto per le equazioni differenziali, c'è però anche un punto dove l'approccio dato qui è meno generale. Infatti, nella Def. 4.1 stiamo chiedendo che $S(t)x$ sia definito per ogni $t \geq 0$ e ogni $x \in \mathcal{X}$ e di fatto escludendo la possibilità che le traiettorie esplodano in tempi finiti (positivi; l'esplosione in tempi finiti negativi è invece ammessa). In realtà, non sarebbe difficile ammettere tale eventualità anche nell'ambito della teoria astratta; tuttavia, occorrerebbe spesso parlare di “tempi ammissibili” e questo porterebbe a riempire definizioni e teoremi di ulteriori dettagli tecnici. Inoltre, segnaliamo che il nostro interesse principale sarà rivolto ai semigruppato cosiddetti “dissipativi”, per i quali la proprietà di esistenza in grande, anche se non fosse postulata, sarebbe comunque conseguenza della condizione di dissipatività.

Se vale la proprietà di unicità all'indietro, allora per ogni $x \in X$, l'insieme $S(t)^{-1}x$ ha al più un elemento. In particolare, se tale insieme non è vuoto, indicheremo con $S(-t)x$ l'unico elemento di $S(t)^{-1}x$. Un modo comodo di rappresentare questa proprietà è dire che, *se due elementi $a, b \in \mathcal{X}$ sono legati dalla relazione $b = S(t)a$ per $t > 0$, allora $S(-t)b$ esiste ed è uguale ad a* . Questa notazione è comoda anche in virtù del fatto che (S2)–(S4) “continuano a valere” anche per i tempi negativi. Più precisamente, abbiamo la

Proposizione 4.9. *Supponiamo che il semigruppso $S(\cdot)$ soddisfi la proprietà di unicità all'indietro. Allora:*

- (a) *Se $S(-t)x$ è definito per qualche $t > 0$, $x \in \mathcal{X}$, allora $S(\tau - t)x$ è definito per ogni $\tau > 0$.*
- (b) *$\forall x \in \mathcal{X}$ e $s, t \in \mathbb{R}$ tali che esista $S(\sigma)x$ per $\sigma := \min\{s, t, s + t\}$, si ha che $S(t)S(s)x = S(s)S(t)x = S(t + s)x$.*
- (c) *Sia $t > 0$ e sia $\{x_n\} \subset \mathcal{X}$ tale che $x_n \rightarrow x$ in $\mathcal{X}$. Supponiamo inoltre che $S(-t)x_n$ sia definita per ogni n e che esista un compatto $K \subset \mathcal{X}$ tale che $S(-t)x_n \in K$ per ogni $n \in \mathbb{N}$. Allora $S(-t)x$ è definita e si ha che $S(-t)x_n \rightarrow S(-t)x$.*
- (d) *$\forall x \in \mathcal{X}$ la (S4) vale per tutti i tempi t nei quali è definita $S(t)x$.*

Prova. Dimostro (a). Posto $v := S(-t)x$ e supponendo $\tau - t < 0$ (altrimenti la proprietà è ovvia), si ha che

$$x = S(t)v = S(t - \tau)S(\tau)v = S(t - \tau)(S(\tau)v), \quad (4.3)$$

da cui, ponendo $z := S(\tau)v$ segue che $S(t - \tau)z = x$, da cui $z = S(\tau - t)x$.

La dimostrazione di (b) si ottiene distinguendo i vari casi possibili sul segno di s e t : è elementare ma noiosa. Si noti che tutti gli oggetti che compaiono nell'enunciato sono ben definiti grazie ad (a).

Dimostriamo (c). Per la compattezza, esistono $\{n_k\}$ ed $y \in \mathcal{X}$ tali che $S(-t)x_{n_k} \rightarrow y$. Ma allora, dato che $t > 0$ e dunque $S(t)$ è continua, si ha che $x_{n_k} = S(t)S(-t)x_{n_k} \rightarrow S(t)y$. Ma poiché $\{x_{n_k}\}$ è una sottosuccessione della $\{x_n\}$, che tende a x , si ha che $S(t)y = x$, da cui $y = S(-t)x$. Infine, per mostrare che l'intera successione $\{S(-t)x_n\}$ tende a $S(-t)x$ si ripete il ragionamento appoggiandosi alla Prop. 5.9.

Per quanto riguarda (d), sia $t \in \mathbb{R}$ tale che è definita $S(t)x$ e sia $\{t_n\}$ una successione di tempi tali che anche $S(t_n)x$ abbia senso e inoltre $t_n \rightarrow t$. Osserviamo che, grazie ad (a), $S(t_n)x = S(t_n - t)S(t)x$. Di qui vediamo che la tesi è banale se $t_n \geq t$ per ogni n . Osservando che nel caso $t_n - t$ cambi segno al variare di n basta distinguere le sottosuccessioni dei $t_n - t$ positivi e negativi, ci resta solo da considerare il caso in cui $t_n \leq t$ per ogni n . Dato allora T tale che $T \leq t_n$ per ogni n e $S(T)x$ esiste, posto $y := S(T)x$, ho subito che $S(t_n)x = S(t_n - T)y \rightarrow S(t - T)y = S(t)x$. ■

Le proprietà delle mappe $S(t)$ permettono di estendere a questo ambiente astratto i concetti di insieme positivamente invariante e di ω -limite. Vediamo anzi di generalizzarli ulteriormente:

Definizione 4.10. *Sia $P \subset \mathcal{X}$. P si dice positivamente invariante se $S(t)P \subset P$ per ogni $t \geq 0$. P si dice negativamente invariante se $P \subset S(t)P$ per ogni $t \geq 0$. Infine, P si dice invariante o completamente invariante se $S(t)P = P$ per ogni $t \geq 0$.*

Come ci si potrebbe aspettare, se vale la proprietà di unicità all'indietro, gli insiemi negativamente invarianti ammettono una caratterizzazione tramite le mappe $S(-t)$, ove $t > 0$.

Proposizione 4.11. *Supponiamo che $S(\cdot)$ verifichi l'unicità all'indietro. Allora, P è negativamente invariante se e solo se $S(-t)x$ è definita per ogni $x \in P$ e $t > 0$ e inoltre $S(-t)P \subset P$ per ogni $t > 0$.*

Prova. Sia innanzitutto P negativamente invariante e sia $x \in P$. Dato $t > 0$, segue allora che $x = S(t)p$ per qualche $p \in P$, da cui $p = S(-t)x$ e dunque $S(-t)x$ esiste e sta in P .

Viceversa, se $S(-t)P \subset P$, allora, dato $x \in P$, ho che $\pi := S(-t)x \in P$, da cui $x = S(t)S(-t)x = S(t)\pi \in S(t)P$. ■

Definizione 4.12. Sia $B \subset \mathcal{X}$. Si chiama ω -limite di B l'insieme

$$\omega\text{-lim } B := \{z \in \mathcal{X} : \exists \{t_n \nearrow\}, \{x_n\} \subset B : S(t_n)x_n \rightarrow z\}. \quad (4.4)$$

Si osservi che, naturalmente, vale l'inclusione

$$\bigcup_{x \in B} \omega\text{-lim } x \subset \omega\text{-lim } B; \quad (4.5)$$

invece, in generale, *non vale* l'inclusione inversa. In un certo senso, l'insieme che contiene i punti limite “di famiglie di traiettorie” può essere *più grande* dell'unione dei punti limite delle singole traiettorie. Se $x \in \mathcal{X}$ e u è la traiettoria del sistema dinamico tale che $u(0) = x$, a volte, con qualche abuso di linguaggio, parleremo di ω -limite di u , anziché di ω -limite di x .

Anche nell'ambito di questa teoria astratta è possibile estendere alcune nozioni relative a punti di equilibrio e stabilità. Cominciamo con la

Definizione 4.13. Sia S un sistema dinamico. Dico che $x \in \mathcal{X}$ è un punto di equilibrio per S se $S(t)x = x$ per ogni $t \geq 0$. Indico con $\mathcal{E}$ l'insieme dei punti di equilibrio di S .

È immediato dimostrare la

Proposizione 4.14. Sia $x \in \mathcal{E}$. Allora $u : \mathbb{R} \rightarrow \mathcal{X}$, $u(t) \equiv x$ per ogni $t \in \mathbb{R}$ è una traiettoria (completa) di S .

Si ha inoltre, come nel caso “concreto”, che

Proposizione 4.15. Sia u una traiettoria di S e supponiamo che $\lim_{s \rightarrow +\infty} u(s) = x \in \mathcal{X}$. Allora $x \in \mathcal{E}$.

Prova. Sia $t > 0$. Per continuità di $S(t)$, si ha che

$$S(t)x = S(t) \lim_{s \rightarrow +\infty} S(s)u(0) = \lim_{s \rightarrow +\infty} S(s+t)u(0) = x. \quad \blacksquare$$

Si noti che l'ipotesi nel precedente enunciato è più forte che chiedere che $\omega\text{-lim } u = \{x\}$. Infatti, se u non ha immagine relativamente compatta in $\mathcal{X}$ potrebbe capitare che x sia l'unico punto limite della traiettoria, ma che x sia raggiunto solo da una sottosuccessione di $u(t)$.

Possiamo anche estendere alcuni concetti che avevamo fin qui introdotto solo nell'ambito dei sistemi *lineari* autonomi.

Definizione 4.16. Sia S un sistema dinamico e sia $x \in \mathcal{E}$. Chiamo varietà stabile associata a x l'insieme $\mathcal{M}_s(x)$ dei punti $z \in X$ tali che esiste una traiettoria completa u di S per cui si abbia $u(0) = z$ e inoltre $u(t) \rightarrow x$ per $t \rightarrow +\infty$. La varietà instabile $\mathcal{M}_i(x)$ di x si definisce in modo analogo, ma sostituendo $-\infty$ a $+\infty$.

Dunque, la varietà stabile di un punto $x \in \mathcal{E}$ è definita come l'unione di tutte le orbite complete generate da traiettorie che tendono a x per t che va a $+\infty$. Si noti infatti che non è limitativo chiedere nella definizione che z sia assunto al tempo 0, dal momento che le traiettorie complete sono invarianti per traslazione. Si noti infine che per ogni $x \in \mathcal{E}$, si ha che $x \in \mathcal{M}_s(x) \cap \mathcal{M}_i(x)$.

Esempio 4.17. Consideriamo l'equazione $y' = 1 - y^2$. Si ha allora $\mathcal{E} = \{-1, 1\}$ e, ad esempio, $\mathcal{M}_s(1) = (-1, 1]$. Infatti, se $z > 1$ è sì vero che la traiettoria $u(t)$ tale che $u(0) = z$ tende a 1 per $t \rightarrow +\infty$. Tuttavia, essa non è completa in quanto esplode in tempi finiti negativi.

Esempio 4.18. Consideriamo un sistema lineare bidimensionale $\mathbf{y}' = A\mathbf{y}$. Si può allora verificare che, in tutti i casi (anche quelli “degeneri”), le varietà stabile $\mathcal{M}_s(0)$ e instabile $\mathcal{M}_i(0)$ coincidono, rispettivamente, con lo spazio stabile E_s e con lo spazio instabile E_i definiti nel precedente capitolo.

Osservazione 4.19. Notiamo che, in quest'ambito più generale, si rinuncia a definire un concetto di “varietà centro”. Si pensi infatti al caso dei sistemi dinamici generati da equazioni non lineari, che costituiscono la principale motivazione di questa teoria astratta. Una definizione di “varietà centro” basata sullo studio degli autovalori della matrice dei coefficienti del sistema linearizzato attorno a un equilibrio x sarebbe poco saggia, perché nel caso vi siano autovalori con parte reale nulla, le traiettorie del sistema possono benissimo allontanarsi dal punto critico x ovvero convergere a questo, e quindi dare luogo a “parti” di varietà stabile o instabile, rispettivamente.

Osservazione 4.20. Usando il concetto di ω -limite si potrebbe fornire alcune estensioni interessanti della precedente definizione. Dato un qualunque sottinsieme $A \subset X$, potremmo dire che z appartiene, ad esempio, alla varietà stabile di A se esiste una traiettoria completa u del sistema tale che $u(0) = z$ e inoltre $\omega\text{-lim } z$ è non vuoto e contenuto in A . Questo è particolarmente interessante in alcuni casi specifici, per esempio quando A è un *ciclo limite*, ossia un'orbita periodica isolata.

Un'ulteriore giustificazione dell'osservazione precedente è data dal seguente teorema, che fornisce un'importante caratterizzazione delle regioni limite dei sistemi bidimensionali. La dimostrazione, che omettiamo, è tecnicamente complicata, ma si basa essenzialmente su considerazioni geometriche almeno intuitivamente elementari (teorema della curva di Jordan) e sul fatto che, naturalmente, due orbite non si possono mai intersecare grazie alla proprietà di unicità. Grazie a questo teorema e alla luce della teoria del prossimo paragrafo risulterà chiaro che fenomeni caotici complicati (si pensi al ben noto, e molto divulgato, fenomeno degli “attrattori strani” - dei quali peraltro non è facile dare una definizione matematicamente rigorosa) non si possono presentare nel caso bidimensionale.

Teorema 4.21. (Poincaré-Bendixson). *Sia S il sistema dinamico associato a un sistema autonomo bidimensionale $\mathbf{y}' = \mathbf{f}(\mathbf{y})$, con $\mathbf{f} \in C^1(\Omega; \mathbb{R}^N)$. Sia $\mathbf{a} \in \Omega$ e sia $K = \omega\text{-lim } \mathbf{a}$. Se K è compatto, non vuoto e non contiene punti di equilibrio, allora K coincide con un'orbita periodica.*

Altre classi importanti di orbite sono introdotte dalla seguente

Definizione 4.22. Sia $\mathcal{O}$ un'orbita completa del sistema dinamico S immagine di una traiettoria u e si ponga $z := u(0)$. $\mathcal{O}$ si dice eteroclina se esistono $x_1, x_2 \in \mathcal{E}$ tali che $z \in \mathcal{M}_s(x_2) \cap \mathcal{M}_i(x_1)$. Se $x_1 = x_2$ e $\mathcal{O}$ non è ridotta a un punto, $\mathcal{O}$ si dice omoclina.

In pratica, le orbite eterocline connettono diversi punti stazionari: nella notazione qui introdotta, $\mathcal{O}$ “parte” da x_1 ed “arriva” in x_2 . Se, invece, abbiamo un'orbita che “parte” e “arriva” in uno stesso punto critico, questa è detta omoclina.

Diamo infine una definizione alternativa di funzione di Liapounov, adatta a questo ambito astratto. Si noti che, da certi punti di vista (ad esempio la regolarità), la definizione è più generale; d'altro canto, la funzione che andiamo a introdurre è definita su tutto lo spazio $\mathcal{X}$ e non solo nell'intorno di un certo punto critico (anche se si potrebbe localizzare con qualche piccolo accorgimento su cui soprassediamo).

Definizione 4.23. Si dice funzione di Liapounov (globale) per il sistema dinamico S una funzione $V \in C^0(\mathcal{X}; \mathbb{R})$ tale che:

- (i) Per ogni $x \in \mathcal{X}$, la $x \mapsto V(S(t)x)$ è non crescente.
- (ii) Se $V(S(t)x) = V(x)$ per qualche $x \in \mathcal{X}$ e $t > 0$, allora $x \in \mathcal{E}$.

Un sistema che ammette una funzione di Liapounov globale si dice sistema gradiente.

Un altro concetto piuttosto interessante (e collegato al precedente) è il seguente:

Definizione 4.24. Si dice integrale primo del sistema dinamico S una funzione $E \in C^0(\mathcal{X}; \mathbb{R})$ tale che $E(S(t)x) = E(x)$ per ogni $x \in \mathcal{X}$ e $t \geq 0$.

Un integrale primo rappresenta dunque una quantità che non varia lungo le traiettorie del sistema dinamico. Nel caso in cui $\mathcal{X}$ è un aperto di $\mathbb{R}^N$, conoscere un integrale primo dà un'informazione importante sul comportamento qualitativo del sistema. Infatti, le orbite sono in questo caso sottoinsiemi delle curve di livello di E , che può essere utile calcolare (in generale, salvo in casi particolarmente fortunati, è solo possibile uno studio locale tramite il teorema delle funzioni implicite).

Esempio 4.25. Particolarmente importante è il caso in cui S è generato da un sistema autonomo bidimensionale della forma

$$\begin{cases} x' = a(x, y) \\ y' = b(x, y) \end{cases} \quad (4.6)$$

ove a e b sono funzioni di classe C^1 . In tal caso, è sempre possibile determinare un integrale primo del sistema, a patto di sapere integrare esplicitamente un'equazione scalare in forma normale del primo ordine. Eliminando infatti formalmente la variabile tempo, ponendo cioè ad esempio

$$\frac{dy}{dx} = \frac{y'(t)}{x'(t)} = \frac{a(x, y)}{b(x, y)}, \quad (4.7)$$

si ha che le curve integrali di (4.7) sono le curve di livello dell'integrale primo. Naturalmente, l'espressione esplicita di E si riesce a ricavare solo quando (4.7) ha una struttura particolare (ad esempio, quando è a variabili separabili).

4.2 Dissipatività e attrattori

La teoria che presentiamo in quest'ultimo paragrafo ha lo scopo di fornire strumenti adatti a descrivere il comportamento delle traiettorie, o meglio di opportune *famiglie di traiettorie* di un sistema dinamico $S(\cdot)$. Lo studio è motivato specialmente dai sistemi non lineari, quando già nel caso bidimensionale esistono traiettorie limitate e non periodiche che non tendono a punti di equilibrio; in dimensione maggiore di 2, poi, possono accadere cose abbastanza complicate e per certi aspetti sorprendenti. In particolare, molto spesso è difficile nei casi concreti riuscire ad avere un'idea qualitativa del comportamento di una soluzione, in altre parole, a “disegnare” la traiettoria corrispondente. Ma il peggio è che talvolta soluzioni che partono da dati iniziali molto vicini (potremmo dire “infinitamente vicini” a patto di precisare il senso di questa locuzione) hanno, nei tempi lunghi, esiti molto diversi.

L'oggetto capace di catturare almeno una parte della dinamica per tempi lunghi di sistemi così caotici è il cosiddetto “attrattore globale”, la cui costruzione sarà portata avanti per passi successivi. Il primo di questi consiste nell'individuare una classe di sistemi, detti *dissipativi*, le cui traiettorie “non divergono” quando t va a infinito. L'introduzione di questa classe motiva in modo naturale la scelta di avere considerato soltanto sistemi definiti per ogni $t \geq 0$, cioè di avere escluso dalla teoria astratta il fenomeno di esplosione in tempi finiti (positivi).

Dal momento che la principale applicazione della teoria astratta è costituita dai sistemi di equazioni differenziali non lineari, per i quali lo spazio delle fasi è in genere un aperto Ω di $\mathbb{R}^N$ (e dunque non è uno spazio metrico completo), vale la pena fare un piccolo sforzo iniziale per includere nella teoria anche questo caso. Pertanto, per tutto il resto del capitolo, $\mathcal{X}$ continuerà ad essere lo spazio metrico “ambiente” su cui è definito il sistema dinamico; tuttavia, ragioneremo spesso su sottoinsiemi $\mathcal{U}$ di $\mathcal{X}$ (in genere positivamente invarianti): si può pensare a $\mathcal{X} = \mathbb{R}^N$ e $\mathcal{U} = \Omega$ nel caso dei sistemi non lineari. In molti casi, anzi, basterebbe che S fosse definito solo su $\mathcal{U}$, ma non faremo questa ipotesi perchè ci converrà in seguito considerare diversi insiemi $\mathcal{U}$. Si noti che non chiederemo nè che $\mathcal{U}$ sia aperto, nè che $\mathcal{X}$ sia completo. La mancanza di completezza verrà superata da richieste di *compattezza* sugli oggetti che introdurremo. I motivi di queste scelte, per ora oscuri, verranno chiariti dal seguito della teoria e dai vari esempi.

Premettiamo una definizione (che in realtà si limita a introdurre un po' di terminologia):

Definizione 4.26. Sia S un sistema dinamico su $\mathcal{X}$ e $\mathcal{U} \subset \mathcal{X}$. Diciamo che un insieme $B \subset \mathcal{X}$ è ben contenuto in $\mathcal{U}$ se la sua chiusura (rispetto alla topologia di $\mathcal{X}$) è contenuta in $\mathcal{U}$.

Dunque, ad esempio, $(0, 1)$ non è ben contenuto in $(-1, 1)$, ma lo è in $(-1, 1]$. In analogia con la definizione precedente, quando parleremo nel seguito di insiemi chiusi, intenderemo sempre rispetto alla topologia di $\mathcal{X}$. In base a questa convenzione, ad esempio, se abbiamo un'equazione differenziale autonoma su $\Omega = (0, +\infty)$ (ad esempio, $y' = \log y$), $(0, 1]$ non è un sottoinsieme chiuso dello spazio delle fasi.

Veniamo ora alla prima definizione “importante”:

Definizione 4.27. Il sistema S si dice dissipativo su $\mathcal{U} \subset \mathcal{X}$ se esiste un insieme limitato $\mathcal{B}_0$ ben contenuto in $\mathcal{U}$ tale che per ogni insieme limitato B ben contenuto in $\mathcal{U}$ esiste $T_B \geq 0$ tale che $S(t)B \subset \mathcal{B}_0$ per ogni $t \geq T_B$. Un tale insieme $\mathcal{B}_0$ si chiama allora insieme assorbente per S su $\mathcal{U}$. Il sistema S si dice (globalmente) dissipativo se si può scegliere $\mathcal{U} = \mathcal{X}$.

Dunque, per definizione, un sistema dinamico è dissipativo quando esso ammette un insieme chiuso e limitato assorbente. La “dissipatività su $\mathcal{U}$ ” non è altro che una localizzazione del concetto globale. La prossima semplice proprietà non dovrebbe risultare sorprendente:

Lemma 4.28. Sia $\mathcal{B}_0 \subset \mathcal{U}$ un insieme compatto assorbente su $\mathcal{U}$ per S . Allora, per ogni B limitato e ben contenuto in $\mathcal{U}$ si ha che $\omega\text{-lim } B$ è un sottoinsieme compatto e non vuoto di $\mathcal{B}_0$.

Prova. Sia $\{a_n\} \subset B$ una successione qualunque e sia $\{t_n \nearrow\}$. Grazie alla Def. 4.27, si ha che $S(t_n)a_n \in \mathcal{B}_0$ almeno per n sufficientemente grande. Ne segue che esiste una sottosuccessione $S(t_{n_k})a_{n_k} \rightarrow z$ ove $z \in \mathcal{B}_0$. Per costruzione, $z \in \omega\text{-lim } B$, che pertanto è non vuoto. Con un ragionamento simile, ma più elementare, si mostra anche che $\omega\text{-lim } B$ è contenuto in $\mathcal{B}_0$ (e dunque è limitato). Infine, resta da vedere che $\omega\text{-lim } B$ è chiuso (e quindi compatto perché contenuto nel compatto $\mathcal{B}_0$). Per questo, mostriamo una caratterizzazione dell’ ω -limite che non fa intervenire le successioni, ma è di tipo puramente topologico. Affermo cioè che

$$\omega\text{-lim } B = \overline{\bigcap_{t \geq 0} \bigcup_{s \geq t} S(s)B}; \quad (4.8)$$

pertanto l’ ω -limite è chiuso perché intersezione di insiemi chiusi. Mostro le due inclusioni che portano alla (4.8). Se z appartiene all’insieme a secondo membro, per ogni scelta di $t = n$ trovo $s_n \geq n$ e $a_n \in B$ tali che

$$|S(s_n)a_n - z| \leq 1/n;$$

inoltre, non è limitativo supporre la $\{s_n\}$ strettamente crescente. Di qui si vede che $z \in \omega\text{-lim } B$.

Viceversa, sia $z \in \omega\text{-lim } B$. Esistono allora $\{t_n \nearrow\}$ e $\{a_n\} \subset B$ tali che $S(t_n)a_n \rightarrow z$. D’altronde, per ogni $t > 0$ si ha che $t_n > t$ definitivamente; e la successione $\{S(t_n)a_n\}$ è dunque definitivamente contenuta in $\bigcup_{s \geq t} S(s)B$. ■

Si noti che il concetto di insieme assorbente può essere ancora interpretato “fisicamente” in termini di “energia”, almeno quando $\mathcal{X}$ è uno spazio vettoriale normato, di norma $|\cdot|$. Infatti, dal momento che ogni insieme che contiene un insieme assorbente è anch’esso assorbente, si può supporre che sia $\mathcal{B}_0 = \overline{B}(0, R_0)$ per qualche $R_0 > 0$ sufficientemente grande. Chiamando “energia” del sistema la funzione quadratica $E(y) := |y|^2/2$, l’insieme assorbente viene dunque a coincidere con le configurazioni di energia inferiore o uguale alla soglia $R_0^2/2$. Secondo questa interpretazione, un sistema dinamico è dissipativo quando, partendo da una famiglia di dati di energia uniformemente inferiore ad una soglia $R^2/2$, questi evolvono fino ad arrivare, dopo un tempo T_R dipendente solo da R , sotto la soglia $R_0^2/2$; fisicamente, il sistema sta “dissipando energia”, cioè la sta per qualche motivo cedendo all’esterno.

Osservazione 4.29. Confrontando l'enunciato del Lemma 4.28 con la Def. 4.27 si noti che c'è un cambio di notazione, apparentemente innocuo, ma di fatto essenziale quando si guarda ad applicazioni infinito-dimensionali della teoria. Nella Definizione, e nell'interpretazione “fisica” in termini di energia, si è chiesto che $\mathcal{B}_0$ fosse *chiuso e limitato*. Invece, l'enunciato del Lemma chiede (e la dimostrazione sfrutta) che $\mathcal{B}_0$ sia *compatto*. A tale proposito, si ricorda che se $\mathcal{X}$ è uno spazio vettoriale normato di dimensione infinita, allora ogni insieme K compatto in $\mathcal{X}$ è anche chiuso e limitato, ma non viceversa. In effetti, nel quadro infinito-dimensionale, il passo difficile per poter applicare questa teoria consiste non nel dimostrare l'esistenza di un insieme assorbente, ma nel far vedere che questo può essere scelto compatto. Torneremo tra breve su questo punto.

Attrattori. Finora abbiamo individuato una classe molto importante di sistemi dinamici, quelli che abbiamo chiamato “dissipativi”, ed abbiamo detto qualcosa sul comportamento all'infinito delle loro traiettorie. Tuttavia, se l'esame dettagliato dei limiti delle singole soluzioni, come abbiamo detto, è in molti casi un compito irrealizzabile, d'altro l'esistenza di un insieme assorbente dice molto poco sul comportamento asintotico del sistema. Per avere maggiori informazioni, ci viene in soccorso l'oggetto chiamato attrattore globale, che ora introdurremo e di cui dimostreremo sotto opportune condizioni l'esistenza. Sarà la struttura dell'attrattore (se si vuole, la sua “forma”) a dirci qualcosa in più sul comportamento per tempi lunghi di S .

Richiamiamo innanzitutto alcuni concetti di distanza tra insiemi. Nel seguito, $d(\cdot, \cdot)$ denota la metrica di $\mathcal{X}$.

Definizione 4.30. Sia A un sottoinsieme di $\mathcal{X}$. Chiamiamo distanza del punto $b \in \mathcal{X}$ dall'insieme A il valore

$$\text{dist}(b, A) := \inf \{d(b, a), a \in A\}, \quad (4.9)$$

ove $d(\cdot, \cdot)$ è la distanza in $\mathcal{X}$. Dati due insiemi $A, B \subset \mathcal{X}$, chiamiamo inoltre distanza unilaterale, o semidistanza o eccesso di B da A il valore

$$\text{dist}(B, A) := \sup \{ \text{dist}(b, A), b \in B \}. \quad (4.10)$$

Chiamiamo infine distanza simmetrica o di Hausdorff degli insiemi A e B

$$\text{dist}_H(B, A) := \max \{ \text{dist}(B, A), \text{dist}(A, B) \}. \quad (4.11)$$

Esercizio 4.31. Assumendo le notazioni della precedente definizione, dimostrare che si ha $\text{dist}(b, A) = 0$ se e solo se $b \in \overline{A}$ e che $\text{dist}(B, A) = 0$ se e solo se $B \subset \overline{A}$.

Parlando molto grossolanamente, la distanza di un punto b dall'insieme A misura la distanza tra b e “il punto di A più vicino a b ” (che non sempre esiste: l'inf nella definizione non è in genere un minimo). L'eccesso di B da A è la distanza (nel senso precedente) da A del “più lontano” (anche qui con pesanti virgolette...) punto di B . Cominciamo a parlare adesso di “attrazione” tra insiemi, cominciando a descrivere il caso in cui si specificano sia l'insieme che “attrae” sia quello i cui sottoinsiemi “vengono attratti”.

Definizione 4.32. Sia S sistema dinamico su $\mathcal{X}$ e $\mathcal{U} \subset \mathcal{X}$. Sia $\mathcal{A}$ un insieme ben contenuto in $\mathcal{U}$. Dico che $\mathcal{A}$ è attraente in $\mathcal{U}$ se per ogni insieme limitato B ben contenuto in $\mathcal{U}$, si ha che

$$\lim_{t \rightarrow \infty} \text{dist}(S(t)B, \mathcal{A}) = 0. \quad (4.12)$$

Dunque, $\mathcal{A}$ è attraente su $\mathcal{U}$, quando le traiettorie che partono da un generico insieme limitato ben contenuto in $\mathcal{U}$ si “avvicinano” a $\mathcal{A}$ (nel senso della distanza tra insiemi) quando t tende a infinito.

Esempio 4.33. Si consideri il sistema dinamico S associato all’equazione $y' = y^3 - y$. È evidente che, su $\mathcal{X} = \mathbb{R}$, non c’è speranza di costruire un insieme attraente; anzi, non siamo neanche in presenza di un semigruppato ai sensi della Def. 4.1, dato che certe traiettorie esplodono in tempi finiti positivi. Per ovviare a questo inconveniente, possiamo “restringere” lo spazio delle fasi, prendendo ad esempio $\mathcal{U} = (-1, 1)$. In questo caso, ci aspettiamo che $\{0\}$ sia un insieme attraente per S . Si noti che $\mathcal{U}$ è limitato e completamente invariante per S . Dunque, non c’è speranza che $\mathcal{U}$ venga attratto da $\{0\}$. Tuttavia, possiamo notare che $\mathcal{U}$ non è “ben contenuto” in se stesso. Invece, se B è ben contenuto in $\mathcal{U}$, è facile verificare che B è attratto da $\{0\}$. Questo è il motivo per cui abbiamo introdotto la classe degli insiemi “ben contenuti” in $\mathcal{U}$.

Osservazione 4.34. Il problema evidenziato nell’esempio precedente è dovuto da un lato alla possibile non completezza di $\mathcal{U}$, dall’altro al concetto di “insieme limitato” che si è voluto usare. Un altro modo di risolvere la situazione (alternativo all’uso degli insiemi “ben contenuti”) è considerare una funzione $f : (-1, 1) \rightarrow \mathbb{R}$ suriettiva e strettamente crescente e definire una nuova distanza su $\mathcal{U} = (-1, 1)$ (riferiamoci per semplicità sempre all’esempio precedente, ove la situazione è particolarmente facile) ponendo $d_1(x, y) := |f(x) - f(y)|$. Tale distanza induce anch’essa la topologia euclidea su $\mathcal{U}$. Tuttavia, rispetto a questa distanza $(-1, 1)$ è completo e non limitato. Inoltre, un sottoinsieme B di $(-1, 1)$ è limitato rispetto a d_1 se e solo se, visto come sottoinsieme di $\mathbb{R}$, è ben contenuto in $(-1, 1)$. Si verifica inoltre facilmente che $\{0\}$ è attraente per S rispetto a d_1 .

La conoscenza di un insieme attraente nel senso della Def. 4.32 è già un’informazione rilevante sul comportamento asintotico di un sistema. Tuttavia, in generale, un generico insieme attraente è ancora “troppo grande” per dare informazioni davvero significative. Consideriamo, infatti, l’equazione $y' = -y$ su $\mathcal{X} = \mathbb{R}$. Evidentemente, ogni sottoinsieme di $\mathbb{R}$ contenente 0 è attraente per il semigruppato associato; tuttavia, l’insieme attraente davvero significativo è quello costituito dal solo 0. Un altro problema è che, di fatto, potrebbe accadere che $\mathcal{A}$ non attragga nulla che sia più grande di se stesso. Ad esempio, dato un generico sistema dinamico S , è ovvio che l’intero spazio $\mathcal{X}$ è attraente per S su $\mathcal{X}$, ma altrettanto ovviamente questa proprietà ha ben poco di significativo.

Ci poniamo, dunque, come prossimo obiettivo, la costruzione di un insieme attraente che sia il più piccolo possibile e che inoltre attragga qualcosa di più grande di se stesso, più precisamente, almeno un suo intorno.

Definizione 4.35. Dico che un sottoinsieme $\mathcal{A}$ di $\mathcal{X}$ chiuso e limitato ha un bacino di attrazione se esiste un insieme aperto $\mathcal{U} \subset \mathcal{X}$ contenente $\mathcal{A}$ tale che $\mathcal{A}$ sia attraente in $\mathcal{U}$.

C'è da osservare che la terminologia usata nella precedente definizione (e anche in quelle che seguiranno) non è universalmente accettata. Anzi, i vari testi sui sistemi dinamici chiamano “insieme attraente” o “attrattore” oggetti di fatto molto diversi. Di fatto, per noi, un insieme attraente $\mathcal{A}$ ha bacino di attrazione se attrae almeno un suo intorno. Si osservi che la richiesta che $\mathcal{A}$ sia chiuso e limitato fatta nella precedente definizione non è di fatto limitativa, dato che vedremo nel seguito (costruendo i cosiddetti “attrattori”) che gli insiemi attraenti “interessanti” sono addirittura compatti. Inoltre, se si dovesse avere un insieme attraente limitato, ma non chiuso, è comunque possibile rimpiazzarlo con la sua chiusura che sarà a maggior ragione attraente.

Il prossimo risultato, che ha una dimostrazione tecnicamente piuttosto complicata, aiuta a caratterizzare il bacino di attrazione di un insieme attraente.

Lemma 4.36. *Supponiamo che nello spazio $\mathcal{X}$ tutti gli insiemi chiusi e limitati siano compatti. Allora, se $\{U_j\}_{j \in J}$ è una qualunque famiglia di aperti tali che $\mathcal{A} \subset U_j$ per ogni $j \in J$ e, inoltre, $\mathcal{A}$ è attraente in U_j , posto $\mathcal{U} := \cup_{j \in J} U_j$, si ha che $\mathcal{A}$ è attraente in $\mathcal{U}$.*

Prova. Sia B un insieme limitato ben contenuto in $\mathcal{U}$. Non è limitativo supporre B chiuso; infatti, altrimenti lo si può sostituire con la sua chiusura, che sarà ancora (ben) contenuta in $\mathcal{U}$. Dimostriamo che $\mathcal{A}$ attrae B . Per ipotesi, B è compatto. Dunque, possiamo estrarre da $\{U_j\}_{j \in J}$ un sottoricoprimento finito di B . Denotiamo per semplicità con $\{U_j\}_{1 \leq j \leq n}$ gli aperti di tale ricoprimento.

Sia, ora, $m \in \mathbb{N}$. Poniamo

$$B_{j,m} := \{x \in B \cap U_j : B(x, 1/m) \subset U_j\}. \quad (4.13)$$

Dal momento che U_j è aperto, è chiaro che per ogni $x \in B \cap U_j$ esiste $m > 0$ sufficientemente grande perché sia $x \in B_{j,m}$. Inoltre, non è difficile mostrare che $B_{j,m}$ è chiuso (e dunque compatto perché sottoinsieme di B) per ogni $j = 1, \dots, n$; $m \in \mathbb{N}$. Infatti, sia $\{x_k\}$ una successione in $B_{j,m}$ convergente a un punto x . Suppongo, per assurdo, che $x \notin B_{j,m}$. In tal caso, esisterebbe $y \in B(x, 1/m) \setminus U_j$. Ma allora, per la disuguaglianza triangolare, si avrebbe che

$$d(y, x_k) \leq d(y, x) + d(x, x_k) < 1/m + d(x, x_k) \quad \forall k \in \mathbb{N}.$$

Ora, dato che $x_k \rightarrow x$, è chiaro che è possibile scegliere k sufficientemente grande perché sia $d(y, x_k) < 1/m$. Ne segue che, per tale scelta di k , y appartiene a $B(x_k, 1/m)$, che è contenuto in U_j perché $x_k \in B_{j,m}$. Di qui l'assurdo; dunque, $B_{j,m}$ è chiuso.

Mostro ora che è possibile scegliere m sufficientemente grande perché sia $B = \cup_{j=1}^n B_{j,m}$. Se ciò non fosse vero, per ogni $m \in \mathbb{N}$ potrei trovare un punto $x_m \in B \setminus \cup_{j=1}^n B_{j,m} = \cap_{j=1}^n (B \setminus B_{j,m})$. Per scelta, dunque, tale punto verifica che

$$\forall j = 1, \dots, n, \quad B(x_m, 1/m) \not\subset U_j. \quad (4.14)$$

Dal momento che x_m varia in B , che è compatto, posso estrarre una sottosuccessione x_{m_h} che converge a un punto $x \in B$ (diverso dall' x di prima!). Ma, essendo $x \in B$, esistono $\bar{m} \in \mathbb{N}$ e $1 \leq \bar{j} \leq n$ tali che $x \in B_{\bar{j}, \bar{m}}$, ossia

$$B(x, 1/\bar{m}) \subset U_{\bar{j}}.$$

Ne segue, allora, che per tutti gli indici h sufficientemente grandi,

$$B(x_{m_h}, 1/(2\bar{m})) \subset U_{\bar{j}};$$

dunque, scelto h in modo tale che sia $m_h > 2\bar{m}$, la (4.14) fornisce l'assurdo.

Mostriamo ora che $\mathcal{A}$ attrae B . Scrivo $B = \cup_{j=1}^n B_{j,m}$ ove m è il valore trovato qui sopra e noto che, per costruzione, $B_{j,m}$ è **ben** contenuto in U_j per ogni $j = 1, \dots, n$ (invece, $B \cap U_j$ non è in genere ben contenuto in U_j); quindi è attratto da $\mathcal{A}$ per ipotesi. Ossia, per ogni $\varepsilon > 0$ e $j = 1, \dots, n$, esiste $T_j > 0$ tale che

$$\text{dist}(S(t)B_{j,m}, \mathcal{A}) \leq \varepsilon \quad \forall t \geq T_j.$$

Posto allora $T := \max_{1 \leq j \leq n} T_j$, è infine evidente che

$$\text{dist}(S(t)B, \mathcal{A}) \leq \varepsilon \quad \forall t \geq T;$$

in altre parole, $\mathcal{A}$ attrae B . ■

Osservazione 4.37. La terminologia “bacino di attrazione” richiede una considerazione aggiuntiva importante. Se valgono le ipotesi del Lemma 4.36, allora è possibile scegliere il bacino di attrazione in modo massimale rispetto all'inclusione. In altre parole, è possibile identificare il *più grande* aperto $\mathcal{U}$ contenente $\mathcal{A}$ e tale che $\mathcal{A}$ attrae i sottoinsiemi ben contenuti in $\mathcal{U}$. Se non vale il Lemma 4.36 (il che accade in quasi tutte le situazioni infinito-dimensionali), questo non è possibile. Cioè, non è possibile costruire un bacino di attrazione “massimale”, a meno di indebolire la proprietà di attrazione. In particolare, se $\mathcal{A}$ è attraente su ciascun insieme di una famiglia di aperti $\{U_j\}_{j \in J}$, posto $\mathcal{U} := \cup_{j \in J} U_j$, allora non è detto che $\mathcal{A}$ sia attraente in $\mathcal{U}$; tuttavia, $\mathcal{A}$ *attrae i punti* di $\mathcal{U}$. Ossia,

$$\forall x \in \mathcal{U}, \quad \lim_{t \rightarrow +\infty} \text{dist}(S(t)x, \mathcal{A}) = 0, \quad (4.15)$$

proprietà evidentemente più debole di quella considerata fin qui.

Il problema evidenziato nella precedente osservazione non si pone quando si è nelle condizioni della seguente

Definizione 4.38. Un sottoinsieme $\mathcal{A}$ chiuso e limitato di $\mathcal{X}$ è globalmente attraente per S se $\mathcal{A}$ è compatto e attraente su $\mathcal{X}$.

In tal caso, infatti, il bacino di attrazione è $\mathcal{X}$ e ogni limitato di $\mathcal{X}$ è attratto. Nel seguito, per semplicità, supporremo sempre di essere in questa situazione (cioè lavoreremo con insiemi globalmente attraenti). Notiamo, tuttavia, che quanto diremo nel seguito, si trasporterà senza difficoltà al caso in cui si voglia considerare un bacino di attrazione (aperto!) $\mathcal{U}$ diverso dall'intero spazio $\mathcal{X}$, ovvero “localizzare” la teoria. L'unica attenzione che va rivolta nell'adattare i prossimi risultati al caso locale è che ogni volta che si prende un insieme limitato, questo deve essere, addizionalmente, *ben contenuto* in $\mathcal{U}$.

Esempio 4.39. Si consideri l'equazione $y' = y - y^3$ su $\mathcal{X} = \mathbb{R}$. Il sistema dinamico associato S è evidentemente dissipativo e si dimostra facilmente che $[-1, 1]$ è attraente per S (anzi, vedremo che è l'attrattore globale). Prendendo però $\mathcal{A} = \{1\}$, si vede subito che $\mathcal{A}$ è localmente attraente per S ed ha bacino di attrazione $\mathcal{U} = (0, +\infty)$. Analogamente, $\mathcal{B} := \{-1\}$ è localmente attraente sul bacino $(-\infty, 0)$.

Esercizio 4.40. Dati un sottoinsieme $A \subset \mathcal{X}$ ed $\varepsilon > 0$, chiamiamo bolla aperta, o intorno aperto, di centro A e raggio ε l'insieme

$$B(A, \varepsilon) := \cup_{b \in A} B(b, \varepsilon). \quad (4.16)$$

Dimostrare allora che, dato $x \in \mathcal{X}$, si ha che

$$x \in B(A, \varepsilon) \iff \text{dist}(x, A) < \varepsilon.$$

Dimostrare inoltre che, se un insieme $\mathcal{B}_0$ è assorbente su $\mathcal{X}$, allora esso è attraente su $\mathcal{X}$. Infine, mostrare che, se $\mathcal{A}$ è un insieme attraente, allora, per ogni $\varepsilon > 0$, $B(\mathcal{A}, \varepsilon)$ è assorbente.

L'esercizio precedente chiarisce parecchie cose sulle relazioni tra insiemi assorbenti e insiemi attraenti. Innanzitutto, se un sistema dinamico ammette un insieme *limitato* attraente $\mathcal{A}$, esso è dissipativo, in quanto ogni intorno chiuso di $\mathcal{A}$ è assorbente. Inoltre, vediamo che la proprietà di "attrattività" è più debole della proprietà di assorbenza. In effetti, se le famiglie di traiettorie che partono da dati iniziali uniformemente limitati tendono a *entrare* dopo un certo tempo in un insieme assorbente, queste tendono solo ad *avvicinarsi* all'insieme attraente. Quale sarà il vantaggio? Mentre l'insieme assorbente è in genere un oggetto molto grande (spesso è una palla), in generale è possibile costruire insiemi attraenti più piccoli. Anzi, più un insieme attraente è piccolo, maggiori sono le informazioni che esso fornisce sul comportamento asintotico delle traiettorie.

Osservazione 4.41. Naturalmente, a volte capita che si riesca a dimostrare che c'è un insieme attraente "molto piccolo", magari addirittura ridotto a un punto, come nell'Esempio 4.33. Questi, in genere, non sono però i casi più interessanti di applicazione della teoria, dal momento che il comportamento asintotico del sistema in questi casi può di solito essere studiato in modo più semplice. Ci sono invece situazioni significative in cui un esame diretto del sistema non fornisce alcuna informazione qualitativa sull'andamento per tempi grandi. È in questi casi che avere un risultato di "piccolezza" o di "caratterizzazione" per un insieme attraente (o, meglio, per l'attrattore globale che definiremo tra breve) può essere utile.

Osservazione 4.42. In generale, nelle applicazioni finito-dimensionali, il problema fondamentale non è tanto stabilire l'esistenza dell'attrattore, quanto investigare le sue proprietà geometriche. Questo, specialmente in dimensione > 2 può essere molto complicato e richiedere strumenti che esulano largamente dalla teoria qui svolta. Diverso è il discorso legato alle applicazioni in dimensione infinita, ove anche l'attrattore può essere infinito-dimensionale (e dunque molto difficile da caratterizzare geometricamente tramite, ad esempio, il comportamento qualitativo delle traiettorie). In questa situazione, già stabilire che l'attrattore è "non troppo grande" (in un senso opportuno) può essere un risultato significativo.

Il prossimo risultato ci aiuterà a identificare *il più piccolo* insieme attraente di un sistema dinamico dissipativo.

Lemma 4.43. *Siano $\mathcal{C}$ un insieme (chiuso) attraente e $\mathcal{B}$ un insieme limitato e completamente invariante per il sistema dinamico $S(\cdot)$. Allora $\mathcal{B} \subset \mathcal{C}$.*

Prova. Grazie alle ipotesi, si ha che $\text{dist}(\mathcal{B}, \mathcal{C}) = \text{dist}(S(t)\mathcal{B}, \mathcal{C}) \rightarrow 0$ per $t \rightarrow \infty$, da cui, grazie all'Es. 4.31 e alla chiusura di $\mathcal{C}$, la tesi. ■

Definizione 4.44. *Sia $\mathcal{A}$ un sottoinsieme di $\mathcal{X}$. Dico che $\mathcal{A}$ è un attrattore (locale) per S se $\mathcal{A}$ è attraente nel bacino di attrazione $\mathcal{U}$, ed è compatto e completamente invariante. Dico che $\mathcal{A}$ è l'attrattore globale o universale del semigruppato se il suo bacino di attrazione è l'intero spazio delle fasi $\mathcal{X}$.*

Si noti che l'esistenza dell'attrattore universale (o quanto meno di un attrattore locale) è un problema non banale, che affronteremo fornendo una caratterizzazione concreta di tale oggetto. Per il momento, osserviamo che, grazie all'Esercizio 4.40, ogni sistema dinamico che ammette l'attrattore universale è dissipativo. Nel seguito ci chiederemo sotto quali condizioni vale anche il viceversa.

Per il momento, cominciamo a notare che, grazie al Lemma 4.43, l'attrattore universale $\mathcal{A}$, se esiste, è unico. Anzi, possiamo dire di più: $\mathcal{A}$ contiene ogni insieme limitato e completamente invariante $\mathcal{B}$ ed è contenuto in ogni insieme chiuso ed attraente $\mathcal{C}$. Dunque, esso è *massimale* nella famiglia degli insiemi limitati e completamente invarianti ed è *minimale* nella famiglia degli insiemi chiusi ed attraenti (ordinate per inclusione).

Osservazione 4.45. L'ipotesi di limitatezza è essenziale nel Lemma 4.43 e nelle osservazioni successive. Considerando ad esempio l'equazione $y' = -y$, è facile provare che l'attrattore globale $\mathcal{A}$ per l'associato sistema dinamico su $\mathcal{X} = \mathbb{R}$ è costituito dal solo $\{0\}$. Tuttavia, $\mathbb{R}$ è un insieme completamente invariante (ma non limitato!) che contiene $\mathcal{A}$.

Veniamo finalmente alla costruzione dell'attrattore universale, presentando il Teorema generale di esistenza, che costituisce il risultato più importante di questo paragrafo:

Teorema 4.46. *Sia $S(\cdot)$ un sistema dinamico dissipativo. Supponiamo anzi che S ammetta un insieme compatto assorbente $\mathcal{B}_0$. Allora S ammette l'attrattore universale $\mathcal{A}$ e questo è dato da*

$$\mathcal{A} = \omega\text{-}\lim \mathcal{B}_0. \quad (4.17)$$

Prova. Grazie al Lemma 4.28, $\mathcal{A} := \omega\text{-}\lim \mathcal{B}_0$ è compatto e non vuoto. In virtù della Def. 4.44, bisogna mostrare che $\mathcal{A}$ è attraente e completamente invariante. Cominciamo dalla prima proprietà.

Procediamo dunque per assurdo e supponiamo che esista un insieme limitato $B \subset \Omega$ che non viene attratto da $\mathcal{A}$. Negare la (4.12) porta allora a dire che esistono $\varepsilon > 0$ ed una successione $\{t_n \nearrow\}$ tali che

$$\text{dist}(S(t_n)B, \mathcal{A}) > \varepsilon \quad \forall n \in \mathbb{N}. \quad (4.18)$$

Ricordando la (4.10), otteniamo allora, $\forall n \in \mathbb{N}$, l'esistenza di un punto $b_n \in B$ tale che $\text{dist}(S(t_n)b_n, \mathcal{A}) > \varepsilon$. D'altronde, poiché $\mathcal{B}_0$ è assorbente e $\{b_n\} \subset B$, segue che esiste $T_B > 0$ tale che per ogni $t \geq T_B$, $n \in \mathbb{N}$, $S(t)b_n \in \mathcal{B}_0$. Ma allora, dato che $\mathcal{B}_0$ è compatto, la successione $\{S(t_n)b_n\}$ ammette un'estratta (che indichiamo per semplicità con la stessa notazione) convergente a un punto z . Infine,

$$z = \lim_{n \rightarrow \infty} S(t_n)b_n = \lim_{n \rightarrow \infty} S(t_n - T_B)S(T_B)b_n \in \omega\text{-}\lim \mathcal{B}_0 = \mathcal{A},$$

poiché $S(T_B)b_n \in \mathcal{B}_0$ e $\{(t_n - T_B) \nearrow\}$, il che dà l'assurdo dal momento che si era costruita la b_n in modo tale che fosse $\text{dist}(S(t_n)b_n, \mathcal{A}) > \varepsilon$ per ogni n .

Dimostriamo infine che $\mathcal{A}$ è completamente invariante. Sia innanzitutto $a \in \mathcal{A}$ e sia $t > 0$. Per costruzione di $\mathcal{A}$, ho allora che

$$a = \lim_{n \rightarrow \infty} S(t_n)b_n, \quad (4.19)$$

ove $\{t_n \nearrow\}$ e $\{b_n\} \subset \mathcal{B}_0$. Ma, allora,

$$S(t)a = S(t)\left(\lim_{n \rightarrow \infty} S(t_n)b_n\right) = \lim_{n \rightarrow \infty} S(t+t_n)b_n \in \mathcal{A},$$

poiché anche $\{(t+t_n) \nearrow\}$. Dunque, $\mathcal{A}$ è positivamente invariante.

Viceversa, proviamo che $\mathcal{A} \subset S(t)\mathcal{A}$. Sia a come in (4.19) un generico elemento di $\mathcal{A}$. Posso allora considerare, almeno per n sufficientemente grande, la successione $\{S(t_n - t)b_n\}$. Poiché $\{(t_n - t) \nearrow\}$ e $\{b_n\} \subset \mathcal{B}_0$ che è limitato, ho che per n sufficientemente grande $S(t_n - t)b_n \in \mathcal{B}_0$, che è compatto. Dunque, posso estrarne una sottosuccessione convergente, corrispondente agli indici $\{n_k\}$, a un punto $z \in \mathcal{A}$. Ma allora

$$S(t)z = S(t)\left(\lim_{k \rightarrow \infty} S(t_{n_k} - t)b_{n_k}\right) = \lim_{k \rightarrow \infty} S(t_{n_k})b_{n_k} = a. \blacksquare$$

Osservazione 4.47. Come si è detto, specialmente in dimensione infinita, se si vuole applicare questa teoria astratta il passo difficile è mostrare che l'assorbente $\mathcal{B}_0$ può essere scelto compatto. In realtà, è possibile indebolire almeno in parte questa condizione ed esistono versioni più generali del Teorema che valgono sotto ipotesi di compattezza più deboli. Ad esempio, una condizione che si riesce talora a verificare è la "compattezza asintotica": si chiede cioè che esista un insieme compatto *attraente*, anziché un compatto assorbente.

Una importante caratterizzazione dell'attrattore universale è data dal seguente

Corollario 4.48. *Sia S un sistema dinamico che ammette l'attrattore universale $\mathcal{A}$. Allora, $\mathcal{A}$ coincide con l'unione di tutte le orbite intere limitate.*

Prova. Dato che ogni orbita intera limitata $\mathcal{O}$ è anche un insieme completamente invariante, $\mathcal{O} \subset \mathcal{A}$ grazie al Lemma 4.43.

Viceversa, sia $a \in \mathcal{A}$. Dal momento che $\mathcal{A}$ è completamente invariante, posso costruire per ricorrenza una successione $\{a_n\} \subset \mathcal{A}$ tale che $a_0 = a$ e $S(1)a_{n+1} = a_n$ per ogni $n \in \mathbb{N}$. A questo punto, pongo $u(t) := S(t+n)a_n$ una volta che $t > -n$. È facile

vedere che la definizione è ben posta (ossia che non c'è dipendenza da n); se infatti $t \geq -m \geq -n$, con $n, m \in \mathbb{N}$, allora si ha che

$$S(t+n)a_n = S(t+m+n-m)a_n = S(t+m)S(n-m)a_n = S(t+m)a_m.$$

Inoltre, si vede facilmente che u è una traiettoria. Infatti, se $\tau \in \mathbb{R}$ e $t > 0$, per $\tau > -n$ abbiamo

$$u(t+\tau) = S(t+\tau+n)a_n = S(t)S(\tau+n)a_n = S(t)u(\tau);$$

dunque, vale la (4.1). L'immagine $\mathcal{O}$ di u è dunque un'orbita intera contenuta in $\mathcal{A}$, è limitata perché $\mathcal{A}$ è compatto, e contiene $a = u(0)$. ■

Osservazione 4.49. Ricordando che, se il bacino di attrazione $\mathcal{U}$ è più piccolo di $\mathcal{X}$, allora $\mathcal{A}$ per definizione attrae solo i sottoinsiemi limitati ben contenuti in $\mathcal{U}$, è facile vedere che la versione “locale” del precedente Corollario si formula dicendo che l'attrattore su un bacino $\mathcal{U}$ coincide con l'unione di tutte le orbite complete, limitate, e ben contenute in $\mathcal{U}$.

Esercizio 4.50. Si consideri il sistema dinamico associato al sistema bidimensionale (in coordinate polari)

$$\begin{cases} \rho' = \rho(1 - \rho) \\ \vartheta' = 1. \end{cases}$$

Prendendo $\mathcal{X} = \mathbb{R}^2$, qual è l'attrattore globale $\mathcal{A}$. Qual è invece l'attrattore nel bacino $\mathcal{U} = \mathbb{R}^2 \setminus \{0\}$? Quali orbite complete vanno, nei due casi, considerate?

L'importanza del Corollario è evidente per le applicazioni: infatti, permette nei casi concreti legati ai sistemi finito-dimensionali descrivere esplicitamente l'attrattore universale in termini delle traiettorie. Occorre comunque ribadire (cfr. Oss. 4.42) che nei sistemi cosiddetti caotici è tutt'altro che semplice individuare quali sono le orbite complete limitate e, soprattutto, tracciarne un andamento qualitativo.

Per concludere, enunciamo una proprietà geometrica importante e naturale dell'attrattore globale $\mathcal{A}$.

Proposizione 4.51. *Supponiamo che il sistema dinamico $S(\cdot)$ ammetta l'attrattore globale $\mathcal{A}$. Allora, se ogni palla di $\mathcal{X}$ è connessa, anche $\mathcal{A}$ è un sottoinsieme connesso di $\mathcal{X}$.*

Prova. Supponiamo, per assurdo, che esistano due aperti disgiunti A_1 e A_2 tali che $\mathcal{A} \subset A_1 \cup A_2$ e $\mathcal{A} \cap A_i$ sia non vuoto per $i = 1, 2$. Poiché $\mathcal{A}$ è limitato, esiste una palla $B = B(x, R)$ di $\mathcal{X}$ contenente $\mathcal{A}$. Ma, allora, poiché $\mathcal{A}$ è invariante,

$$\mathcal{A} = S(t)\mathcal{A} \subset S(t)B \quad \forall t > 0.$$

Dato che $S(t)$ è continua, $S(t)B$ è connesso e dunque, prendendo $t = n$, esiste sicuramente

$$x_n \in S(n)B \setminus (A_1 \cup A_2). \quad (4.20)$$

Ma, poiché $\mathcal{A}$ è compatto, per ogni n esiste $a_n \in \mathcal{A}$ che realizza il minimo della distanza da x_n , ossia

$$0 \leq d(x_n, a_n) = d(x_n, \mathcal{A}) \leq \text{dist}(S(n)B, \mathcal{A}). \quad (4.21)$$

Usando ancora la compattezza di $\mathcal{A}$ estraggo una sottosuccessione a_{n_k} che tende ad $a \in \mathcal{A}$. Poiché $\mathcal{A}$ è attraente, passando al limite per $k \rightarrow \infty$ nella (4.21) scritta per $\{n_k\}$, ottengo allora che $x_{n_k} \rightarrow a$, che dà l'assurdo a causa della (4.20).

Esempio 4.52. Si consideri ancora l'equazione $y' = y - y^3$ sullo spazio delle fasi (completo!) $\mathcal{X} = (-\infty, -1] \cup [1, +\infty)$. Si verifica facilmente che allora, l'attrattore $\mathcal{A}$ è dato dall'insieme sconnesso $\{-1, 1\}$.

5 Il teorema di Peano

5.1 Spazi metrici, completezza, compattezza

Nel seguito, (X, d) o, più semplicemente, X sarà uno *spazio metrico* con distanza d . Supponiamo in particolare che siano già note (si trovano comunque sui libri di testo del primo anno) le definizioni di distanza, spazio metrico, successione di Cauchy, successione convergente. Ricordiamo che X si dice *completo* se ogni successione di Cauchy in X è convergente. Dato che ogni sottoinsieme A di X è a sua volta uno spazio metrico, ha senso parlare di completezza di A . Ricordiamo (vedi anche il Paragrafo 1.1 di queste dispense) che, se X è completo, allora A è completo se e solo se è chiuso. Se A è completo e X no, A è comunque chiuso in X . Ricordiamo infine che uno spazio X si dice *compatto per successioni* se ogni successione a valori in X ammette un'estratta convergente. Se si considera un sottoinsieme $A \subset X$, dire che A è compatto per successioni richiede che l'estratta converga a un elemento di A . Se è invece dato $A \subset X$ tale che ogni successione in A ammette un'estratta convergente, ma non necessariamente a un elemento di A , diciamo che A è *relativamente compatto (per successioni)*. È facile dimostrare (lo si faccia per esercizio) che A è relativamente compatto per successioni se e solo se la sua chiusura $\overline{A}$ è compatta per successioni.

Un primo rapporto tra compattezza e completezza è dato dal seguente

Teorema 5.1. *Se X è compatto per successioni, allora è completo.*

Prova. Sia $\{x_n\} \subset X$ di Cauchy e sia $\{x_{n_k}\}$ un'estratta convergente a un punto x . Fissato $\varepsilon > 0$ esistono (per la condizione di Cauchy e la convergenza di $\{x_{n_k}\}$, rispettivamente) n_1 e n_2 tali che

$$d(x_n, x_m) \leq \varepsilon \quad \forall n, m \geq n_1, \quad (5.1)$$

$$d(x_{n_k}, x) \leq \varepsilon \quad \forall k \geq n_2. \quad (5.2)$$

Prendendo $m = n_k$ in (5.1), notando che $n_k \geq k$ e usando la disuguaglianza triangolare, la tesi segue. ■

Accanto alla compattezza per successioni dovrebbe essere noto il concetto topologico di compattanza: se X è uno spazio metrico (o topologico), X è compatto (in senso topologico) se da ogni suo ricoprimento $\{U_j\}_{j \in J}$ composto da insiemi aperti, dove J è un insieme di indici di qualunque cardinalità, è possibile estrarre un sottoricoprimento finito. Per gli spazi metrici, i due concetti di compattezza (per successioni e in senso topologico) sono equivalenti, come mostra l'importante teorema che enunceremo fra un attimo. Ciò invece non vale, anche se i controesempi non sono del tutto banali, in spazi topologici più generali. Cominciamo col dare una

Definizione 5.2. Se $A \subset X$ si chiama *diametro di A* il numero

$$\text{diam}(A) := \sup\{d(x, y) : x, y \in A\}.$$

Un insieme $A \subset X$ si dice *totalmente limitato* se per ogni $\varepsilon > 0$ esiste una famiglia finita $\{A_i\}_{1 \leq i \leq n}$ di sottoinsiemi di X di diametro minore o uguale a ε tale che $A \subset \bigcup_{1 \leq i \leq n} A_i$.

Si noti che non sarebbe limitativo chiedere nella definizione che sia $A_i \subset A$ (in caso contrario, basterebbe infatti rimpiazzare A_i con $A_i \cap A$). Allo stesso modo, si potrebbe supporre che gli insiemi A_i siano palle di raggio ε (o, meglio $\varepsilon/2$, se vogliamo che il diametro sia effettivamente ε). Infine, poichè il diametro di un insieme non varia passando alla chiusura, possiamo anche supporre che gli A_i siano insiemi chiusi. Dunque, in sostanza, un insieme A è totalmente limitato quando è possibile ricoprirlo con un numero finito n di insiemi di diametro arbitrariamente piccolo. Ovviamente, tale n dipenderà dalla scelta di ε . È anche immediato dimostrare che in $\mathbb{R}$, o in $\mathbb{R}^N$, un insieme è limitato se e solo se è totalmente limitato. Invece, come vedremo più avanti, in situazioni più generali (ancora una volta quelle in qualche modo connesse a spazi di dimensione *infinita* – sebbene non sia ovvio dire che cosa è la “dimensione” di uno spazio metrico) questo fatto cessa di essere vero. Veniamo ora al teorema.

Teorema 5.3. *Le seguenti condizioni sono equivalenti:*

- (a) X è compatto per successioni.
- (b) X è completo e totalmente limitato.
- (c) X è compatto in senso topologico.

Prova. (a) $\Rightarrow$ (b). Grazie al risultato precedente, basta mostrare che X è totalmente limitato. Se ciò non fosse vero, esisterebbe $\varepsilon > 0$ tale che X non si potrebbe ricoprire con un numero finito di palle di raggio ε . In particolare, dato un qualunque punto $x_1 \in X$, $B_1 := \overline{B}(x_1, \varepsilon)$ non ricopre X . Dunque, esiste $x_2 \in X \setminus B_1$. Considerando $B_2 := \overline{B}(x_2, \varepsilon)$, B_1 e B_2 non ricoprono X . Iterando il procedimento, costruisco una successione $\{x_n\}$ tale che $d(x_i, x_j) > \varepsilon$ per ogni $i \neq j \in \mathbb{N}$. Da tale successione non si può estrarre alcuna sottosuccessione convergente; di qui l'assurdo.

(b) $\Rightarrow$ (c). Sia $\{U_j\}_{j \in J}$ un ricoprimento aperto di X e supponiamo che non se ne possa estrarre alcun sottoricoprimento finito. Poiché X è totalmente limitato, esiste una famiglia finita $\mathcal{F}_1$ di insiemi *chiusi*, di diametro minore di 1 e tali che X è l'unione degli insiemi di $\mathcal{F}_1$. Per la negazione della tesi, è chiaro che esiste un elemento X_1 di $\mathcal{F}_1$ che non può essere ricoperto da un numero finito di insiemi del ricoprimento $\{U_j\}$. Sia esso X_1 . Poiché anche X_1 è totalmente limitato, posso iterare il procedimento scrivendo X_1 come unione di una famiglia finita $\mathcal{F}_2$ di suoi sottoinsiemi chiusi di diametro minore di $1/2$. Uno di questi, detto ad esempio X_2 , non potrà essere ricoperto da un numero finito degli $\{U_j\}$.

In questo modo, posso ottenere una famiglia $\{X_n\}_{n \in \mathbb{N}}$ di sottoinsiemi chiusi di X tali che $X_{n+1} \subset X_n$ per ogni n , $\text{diam}(X_n) \leq 1/n$, nessuno dei quali può essere ricoperto da un numero finito degli $\{U_j\}$. Preso, per ogni n , un qualunque punto $x_n \in X_n$, è facile verificare che la successione $\{x_n\}$ è di Cauchy; dunque, per l'ipotesi di completezza, essa ammette un limite x . Poiché la successione è definitivamente

contenuta in ciascuno degli X_n e questi sono chiusi, segue che $x \in X_n$ per ogni n . D'altro canto, poiché $\{U_j\}$ è un ricoprimento aperto di X , esisteranno $\varepsilon > 0$ e $\bar{j} \in J$ tali che $B(x, \varepsilon) \subset U_{\bar{j}}$. Ma, allora, scelto n in modo tale che sia $1/n < \varepsilon$, essendo $x \in X_n$, si ha che

$$d(x, y) \leq \text{diam } X_n \leq 1/n < \varepsilon \quad \forall y \in X_n,$$

da cui $X_n \subset B(x, \varepsilon) \subset U_{\bar{j}}$ per tale n , il che è assurdo.

(c) $\Rightarrow$ (a). Sia, per assurdo, $\{x_n\}$ una successione in X da cui non si può estrarre alcuna sottosuccessione convergente. Segue allora che ogni punto $x \in X$ ammette un intorno (aperto) U_x che contiene solo un numero finito di punti della successione. In questo modo, si costruisce un ricoprimento $\{U_x\}_{x \in X}$ dal quale, evidentemente, non è possibile estrarre alcun sottoricoprimento finito. ■

Notiamo che dal teorema segue in particolare che, dato un sottoinsieme A di uno spazio metrico completo X , A è totalmente limitato se e solo se ha chiusura compatta (o compatta per successioni). In particolare, se $\{a_n\}$ è una successione a valori in un insieme totalmente limitato A , essa ammette una sottosuccessione convergente (ma non necessariamente a un elemento di A).

5.2 Il teorema di Ascoli

In questo paragrafo vogliamo introdurre lo strumento fondamentale che sarà usato nella dimostrazione del teorema di esistenza di Peano. Tale risultato, chiamato teorema di Ascoli, caratterizza gli insiemi (relativamente) compatti di funzioni continue rispetto alla norma $\|\cdot\|_\infty$ del massimo e giocherà grosso modo il ruolo che aveva il teorema delle contrazioni nella prova del teorema di esistenza e unicità in piccolo. Per dare l'enunciato, abbiamo però bisogno di un certo numero di definizioni preliminari, anche perché vogliamo presentare una versione il più possibile generale del teorema. Lo sforzo aggiuntivo sarà compensato dal fatto che l'enunciato che forniremo potrà essere applicato (non nel presente contesto, però) anche a molte altre situazioni importanti nelle applicazioni (ad esempio nello studio di equazioni a derivate parziali, in particolare di tipo evolutivo).

Esempio 5.4. Consideriamo lo spazio $\mathcal{X} = C([0, 1]; \mathbb{R})$ dotato della norma del massimo. Si è già visto che $\mathcal{X}$ è un Banach e dunque, in particolare, uno spazio metrico completo. Vogliamo vedere che i sottoinsiemi chiusi e limitati di $\mathcal{X}$ non sono necessariamente compatti. Grazie ai risultati del paragrafo precedente, è sufficiente esibire una successione limitata a valori in $\mathcal{X}$ dalla quale non sia possibile estrarre alcuna sottosuccessione convergente. Prendiamo ad esempio $\{f_n\}$ data da

$$f_n(t) = \begin{cases} nt & \text{se } t \in [0, 1/n] \\ 1 & \text{se } t \in (1/n, 1]. \end{cases}$$

È evidente che f_n è limitata in norma e converge puntualmente alla funzione limite f identicamente uguale a 1 salvo in 0 dove vale 0. Ciò implica che da $\{f_n\}$ non si può estrarre alcuna sottosuccessione uniformemente convergente; infatti, per l'unicità del limite, questa dovrebbe tendere a f , che è discontinua.

Dunque, per avere la (relativa) compattezza di un insieme limitato in norma di funzioni continue è necessaria una condizione aggiuntiva. Per introdurla, mettiamoci in un contesto più generale del precedente e consideriamo una coppia di spazi metrici X, Z , con X compatto e Z completo. Allora, è facile verificare che, innanzitutto,

$$d(f, g) := \max\{d_Z(f(x), g(x)) : x \in X\} \quad (5.3)$$

è una metrica sullo spazio $C(X; Z)$; inoltre, tale spazio è completo. La dimostrazione di tali affermazioni si svolge con poche modifiche rispetto al caso noto in cui X e Z sono sottoinsiemi di $\mathbb{R}^N$. Si noti soltanto che in questo contesto Z potrebbe non avere una struttura vettoriale (e, nel caso, neanche $C(X; Z)$ la avrebbe): dunque, va usata con qualche cura la disuguaglianza triangolare. Si noti anche che la completezza di Z è essenziale per avere completezza di $C(X; Z)$. Nella (5.3), e nel seguito, sto indicando con d_X e d_Z le distanze in X e in Z , rispettivamente, e con d la distanza in $C(X; Z)$.

La seguente proprietà è alla base della caratterizzazione dei compatti di $C(X; Z)$:

Definizione 5.5. *Sia $\mathcal{F} \subset C(X; Z)$. La famiglia $\mathcal{F}$ si dice equicontinua in un punto $x \in X$ se per ogni $\varepsilon > 0$ esiste un intorno $U_x \in \mathcal{U}(x)$ tale che $d_Z(f(x), f(y)) \leq \varepsilon$ per ogni $y \in U_x$ e per ogni $f \in \mathcal{F}$.*

Naturalmente, se si vuole, non è limitativo supporre che l'intorno U_x sia della forma $B(x, \delta)$ per $\delta > 0$ dipendente da ε . La cosa importante, però, è che δ non deve dipendere dall'elemento $f \in \mathcal{F}$. A priori, invece, ammettiamo che δ possa dipendere da x , anche se vedremo che, grazie alla compattezza di X , ciò è di fatto inessenziale. Tornando al precedente Esempio 5.4, osserviamo che la condizione di equicontinuità non è verificata dalla $\mathcal{F} = \{f_n\}$.

Prima di enunciare il teorema, occorre un'ultima osservazione. Se nel caso in cui $Z = \mathbb{R}$ è sufficiente supporre la limitatezza in norma delle funzioni considerate, passando a uno spazio Z più generale può sorgere il problema che i limitati di Z in genere non sono relativamente compatti. Dunque, la limitatezza in norma dovrà essere sostituita nell'enunciato da una condizione (più forte) di “puntuale compattezza delle immagini” (vedi (b) qui sotto). Anche qui ammetteremo per adesso che non vi sia uniformità in $x \in X$ (ma anche in questo caso la cosa sarà di fatto automatica, vedi il successivo Corollario 5.7). Veniamo dunque al teorema.

Teorema 5.6. (Teorema di Ascoli). *Sia $\mathcal{F} \subset C(X; Z)$. Allora $\mathcal{F}$ è relativamente compatto se e solo se:*

- (a) *per ogni $x \in X$, $\mathcal{F}$ è equicontinuo in x ;*
- (b) *$\mathcal{F}$ è puntualmente relativamente compatto. Ovvero, per ogni $x \in X$, $\mathcal{F}(x) := \cup_{f \in \mathcal{F}} \{f(x)\}$ è relativamente compatto in Z .*

Prova. (1). Sia $\mathcal{F}$ relativamente compatto e dimostriamo (a) e (b). Innanzitutto, data $\{f_n\} \subset \mathcal{F}$, per ipotesi questa ammette una $\{f_{n_k}\}$ convergente a un limite $f \in C(X; Z)$. Grazie alla continuità dell'operatore di valutazione $V_x : C(X; Z) \rightarrow Z$, $V_x : f \mapsto f(x)$ (che segue dalla scelta della norma su $C(X; Z)$), abbiamo allora che $f_{n_k}(x) \rightarrow f(x)$ in Z per ogni $x \in X$; ovvero, vale (b). Per quanto riguarda (a), osserviamo che, per il Teorema 5.3, $\mathcal{F}$ è totalmente limitato. Scelto arbitrariamente

$\varepsilon > 0$ ricopriamo dunque $\mathcal{F}$ con un numero finito n di insiemi F_i ciascuno dei quali di diametro $< 2\varepsilon$. Si può, anzi, supporre $F_i = B(f_i, \varepsilon)$ per qualche $f_i \in \mathcal{F}$. Osserviamo ora che, per il teorema di Heine (che si estende anch'esso senza difficoltà a questo contesto generale – verificare per esercizio), ciascuna delle f_i è uniformemente continua. Dunque, in corrispondenza di ε , esistono numeri $\delta_i > 0$, per $i = 1, \dots, n$, tali che

$$d_Z(f_i(x), f_i(y)) \leq \varepsilon \quad \forall x, y \in X : d_X(x, y) < \delta_i.$$

Si ponga ora $\delta := \min\{\delta_i : i = 1, \dots, n\}$. Data una qualunque $f \in \mathcal{F}$, sia $\bar{i}$ l'indice i tale che $f \in F_i$. Si ha allora che

$$\begin{aligned} d_Z(f(x), f(y)) &\leq d_Z(f(x), f_{\bar{i}}(x)) + d_Z(f_{\bar{i}}(x), f_{\bar{i}}(y)) + d_Z(f_{\bar{i}}(y), f(y)) \\ &\leq 3\varepsilon, \quad \forall x, y \in X : d_X(x, y) < \delta. \end{aligned} \quad (5.4)$$

Questa è proprio la tesi; anzi, abbiamo dimostrato qualcosa di leggermente più forte, ossia una proprietà di *uniforme* (in x) equicontinuità; vale a dire che δ può essere scelto indipendente da $x \in X$ oltreché da $f \in \mathcal{F}$.

(2). Supponiamo ora che valgano (a) e (b) e dimostriamo che $\mathcal{F}$ è totalmente limitato. Sia dunque $\varepsilon > 0$. Innanzitutto, grazie ad (a), per ogni $x \in X$ possiamo trovare un aperto $U_x \in \mathcal{U}(x)$ buono per la condizione di equicontinuità, ossia tale che

$$d_Z(f(x), f(y)) \leq \varepsilon \quad \forall y \in U_x, \quad \forall f \in \mathcal{F}.$$

È chiaro che $\{U_x\}_{x \in X}$ è un ricoprimento aperto di X . Estraiamo allora un sottoricoprimento finito, che ribattezziamo $\{U_i\}_{i=1, \dots, n}$. Sia x_i il punto $x \in X$ tale che $U_i = U_{x_i}$.

Poniamo ora, per $i = 1, \dots, n$, $Z_i := \overline{\mathcal{F}(x_i)}$. Grazie a (b), Z_i è compatto in Z . Introduciamo l'applicazione $\mathcal{P} : C(X; Z) \rightarrow Z^n$, definita da

$$\mathcal{P}(f) := (f(x_1), f(x_2), \dots, f(x_n)).$$

Per costruzione, si ha che l'immagine di $\mathcal{F}$ tramite $\mathcal{P}$ è contenuta in $Z_1 \times Z_2 \times \dots \times Z_n$. Poiché Z_i è compatto in Z per ogni i , $Z_1 \times Z_2 \times \dots \times Z_n$ è compatto (come sottoinsieme di Z^n) grazie al teorema di Tychonoff (che, ricordiamo, afferma che il prodotto topologico di spazi compatti è compatto: peraltro qui la proprietà è banale perché stiamo prendendo il prodotto di un numero *finito* di spazi). Sullo spazio Z^n stiamo naturalmente considerando la *metrica prodotto*

$$d_{Z^n}(\mathbf{z}^1, \mathbf{z}^2) := \sum_{i=1, \dots, n} d_Z(z_i^1, z_i^2),$$

ove $\mathbf{z}^1$ e $\mathbf{z}^2$ sono vettori di Z^n di componenti z_i^1 e z_i^2 , rispettivamente. È facile vedere che tale metrica in effetti induce proprio la topologia prodotto su Z^n . Abbiamo, di conseguenza, che $\mathcal{P}(\mathcal{F})$ è relativamente compatto poiché contenuto nel compatto $Z_1 \times Z_2 \times \dots \times Z_n$. Dunque, è possibile ricoprire $\mathcal{P}(\mathcal{F})$ con un numero finito m di insiemi di diametro $< 2\varepsilon$ rispetto alla metrica prodotto, i quali possono essere ancora una volta supposti della forma $B(\mathcal{P}(f_j), \varepsilon)$ per $f_j \in \mathcal{F}$, $j = 1, \dots, m$.

A questo punto, possiamo dimostrare che $\mathcal{F}$ è totalmente limitato. In particolare, vediamo che $\mathcal{F}$ si ricopre con la famiglia $\{F_j\}$, ove $F_j := B(f_j, 3\varepsilon)$, $j = 1, \dots, m$. Siano

infatti dati un punto qualunque $x \in X$ e un elemento qualunque $f \in \mathcal{F}$. Sia $\bar{i}$ l'indice i tale che $x \in U_i$ e sia $\bar{j}$ l'indice j per cui $\mathcal{P}(f) \in B(\mathcal{P}(f_j), \varepsilon)$. La seconda proprietà implica che

$$d_{Z^n}((f(x_1), \dots, f(x_n)), (f_{\bar{j}}(x_1), \dots, f_{\bar{j}}(x_n))) \leq \varepsilon,$$

da cui, in particolare, $d_Z(f(x_{\bar{i}}), f_{\bar{j}}(x_{\bar{i}})) \leq \varepsilon$.

A questo punto, abbiamo che

$$d_Z(f(x), f_{\bar{j}}(x)) \leq d_Z(f(x), f(x_{\bar{i}})) + d_Z(f(x_{\bar{i}}), f_{\bar{j}}(x_{\bar{i}})) + d_Z(f_{\bar{j}}(x_{\bar{i}}), f_{\bar{j}}(x)) \leq 3\varepsilon,$$

dove il primo e il terzo termine sono stati maggiorati grazie all'equicontinuità. Poiché tale disuguaglianza vale per ogni $x \in X$, prendendo il massimo rispetto a x , otteniamo $d(f, f_{\bar{j}}) \leq 3\varepsilon$, ossia $f \in B(f_{\bar{j}}, 3\varepsilon)$. ■

Il prossimo risultato mostra, come promesso, che dalla proprietà locale (b) segue, nelle ipotesi del teorema, una compattezza globale dell'immagine di $\mathcal{F}$. Naturalmente, perché ciò sia vero, è essenziale che valga anche l'equicontinuità (a).

Corollario 5.7. *Sotto le ipotesi del Teorema 5.6, l'insieme $\mathcal{F}(X) := \{f(x) : f \in \mathcal{F}, x \in X\}$ è relativamente compatto in Z .*

Prova. Mostro che $\mathcal{F}(X)$ è relativamente compatto per successioni. Sia dunque $\{z_n\} \subset \mathcal{F}(X)$ data, per $n \in \mathbb{N}$, da $z_n = f_n(x_n)$ ove $f_n \in \mathcal{F}$ e $x_n \in X$. Poiché X è compatto, posso estrarre una $\{x_{n_k}\}$ convergente a $x \in X$. Inoltre, grazie al teorema di Ascoli, posso estrarre (come al solito non in modo indipendente, ma a cascata rispetto alla precedente estrazione) una $\{f_{n_{k_h}}\}$ convergente in $C(X; Z)$ a un limite f non necessariamente appartenente a $\mathcal{F}$. Pongo $z := f(x)$. Ho allora che

$$d_Z(z_{n_{k_h}}, z) = d_Z(f_{n_{k_h}}(x_{n_{k_h}}), f(x)) \leq d_Z(f_{n_{k_h}}(x_{n_{k_h}}), f_{n_{k_h}}(x)) + d_Z(f_{n_{k_h}}(x), f(x)).$$

Poiché $\mathcal{F}$ è equicontinuo in x , dato $\varepsilon > 0$, trovo $\delta > 0$ tale che $d_Z(g(x), g(y)) < \varepsilon$ per ogni $y \in B(x, \delta)$ e ogni $g \in \mathcal{F}$. D'altro canto, per h sufficientemente grande, in modo dipendente da δ e conseguentemente da ε , ho che $x_{n_{k_h}} \in B(x, \delta)$. Grazie a tali considerazioni, segue subito dalla formula precedente che, per h sufficientemente grande, $d_Z(z_{n_{k_h}}, z) < 2\varepsilon$. ■

Dunque, mettendo assieme teorema e corollario, vediamo che se si hanno contemporaneamente equicontinuità e relativa compattezza delle immagini puntualmente in $x \in X$, allora tali proprietà valgono in senso uniforme in x .

Notiamo anche che, nel caso in cui X è un sottoinsieme di $\mathbb{R}^N$ e $Z = \mathbb{R}^M$, allora l'enunciato del teorema di Ascoli si semplifica. Si ha cioè che un sottoinsieme $\mathcal{F}$ di $C(X; \mathbb{R}^M)$ è relativamente compatto se e solo se è (puntualmente o uniformemente in $x \in X$) equicontinuo e limitato.

Esercizio 5.8. Sia $\mathcal{F} \subset C^1(I; \mathbb{R})$ dove I è un intervallo chiuso e limitato di $\mathbb{R}$ e lo spazio $C^1(I; \mathbb{R})$ è dotato della norma

$$\|f\|_{C^1} := \|f\|_{\infty} + \|f'\|_{\infty}.$$

Mostrare innanzitutto che tale norma rende $C^1(I; \mathbb{R})$ uno spazio di Banach. Inoltre, usando il teorema di Ascoli, provare che se $\mathcal{F}$ è limitato rispetto alla norma di C^1 , allora è relativamente compatto in $C(I; \mathbb{R})$.

Concludiamo questo paragrafo dimostrando un risultato talvolta utile per verificare la convergenza di una successione a valori in $\mathcal{X}$. Osserviamo che la dimostrazione non fa uso delle proprietà della metrica; in effetti, il risultato continua a valere se $\mathcal{X}$ è un qualunque spazio topologico di Hausdorff.

Proposizione 5.9. *Sia $\mathcal{X}$ uno spazio metrico e $\{x_n\} \subset \mathcal{X}$ una successione. Supponiamo che esista $x \in \mathcal{X}$ tale che ogni sottosuccessione $\{x_{n_k}\}$ della $\{x_n\}$ ammetta un'ulteriore estratta $\{x_{n_{k_h}}\}$ convergente a x . Allora $x_n \rightarrow x$*

Prova. Notiamo innanzitutto che l'enunciato richiede esplicitamente che il limite x sia indipendente dalla scelta della $\{x_{n_k}\}$. Per assurdo, supponiamo che x_n non tenda a x . Allora esistono $U \in \mathcal{U}(x)$ e $\{x_{n_k}\}$ estratta da $\{x_n\}$ tali che per ogni $k \in \mathbb{N}$, $x_{n_k} \notin U$. Ma, per ipotesi, esiste $\{x_{n_{k_h}}\}$ estratta da $\{x_{n_k}\}$ e convergente a x , il che dà l'assurdo perché allora dovrebbe essere $x_{n_{k_h}} \in U$ almeno per h sufficientemente grande. ■

5.3 Esistenza senza unicità

Torniamo a considerare il Problema di Cauchy (1.16), che riscriviamo:

$$\begin{cases} \mathbf{y}' = \mathbf{f}(t, \mathbf{y}) \\ \mathbf{y}(t_0) = \mathbf{y}_0. \end{cases} \quad (5.5)$$

Per il momento, $\mathbf{f}$ sarà, come nel Cap. 1, una funzione definita e *continua* su un aperto qualunque $\Omega \subset \mathbb{R}^{N+1}$ a valori in $\mathbb{R}^N$. Come annunciato, non supporremo alcuna ipotesi di Lipschitz su $\mathbf{f}$. Il nostro obbiettivo in questo paragrafo è ripercorrere la teoria generale cercando di capire che cosa si salva di questa e che cosa viene meno senza la Lipschitzianità di $\mathbf{f}$.

Il risultato principale sarà l'esistenza di almeno una soluzione in piccolo, ed essenzialmente a cadere sarà la proprietà di unicità. Per dimostrare questo, cominciamo col notare che continua a valere il Lemma 1.16, ovvero il Problema (5.5) ammette la riformulazione come equazione integrale di Volterra. Il teorema di “esistenza senza unicità”, dovuto a Peano, sarà dimostrato con un approccio diverso rispetto al caso Lipschitz: infatti, dapprima si tratterà il caso ove $\mathbf{f}$ è limitata e definita su una “striscia verticale” (come nel risultato di esistenza in grande). Dal risultato “globale” che si proverà sotto queste ipotesi aggiuntive, si passerà poi al teorema vero e proprio tramite una procedura di “localizzazione”. In quanto segue ci occuperemo, per semplicità, del solo problema “in avanti”, quello all'indietro potendosi trattare con tecniche analoghe. Come al solito, è inoltre possibile “incollare” le soluzioni in avanti e quelle all'indietro, ottenendo soluzioni bilaterali. Notiamo anche che $\mathbf{f}$ si supporrà inizialmente definita su una striscia verticale *chiusa*, ossia si considereranno tempi $t \in [t_0, t_0 + \tau]$, ove $\tau > 0$. Questo piccolo strappo rispetto alla convenzione per cui si cercano di solito soluzioni definite su intervalli aperti (o, nel caso del problema in avanti, semiaperti a destra) è motivato dalla necessità di usare il Teorema di Ascoli e, comunque, sarà ricucito nei risultati successivi. Ricordiamo ora una proprietà, la cui dimostrazione è pressoché immediata:

Lemma 5.10. Sia $I \subset \mathbb{R}$ un intervallo chiuso e sia $\phi : I \rightarrow \mathbb{R}^N$ una funzione regolare a tratti (ossia continua in I e derivabile salvo in un insieme finito $I_0 \subset I$). Sia, inoltre, $M > 0$ tale che $|\phi'(t)| \leq M$ per ogni $t \in I \setminus I_0$. Allora ϕ è Lipschitziana di costante $\sqrt{N}M$.

Veniamo finalmente all'enunciato del Teorema:

Teorema 5.11. Sia $\tau > 0$, $\mathbf{f} \in C^0([t_0, t_0 + \tau] \times \mathbb{R}^N; \mathbb{R}^N)$ limitata in modulo da $M > 0$. Allora, il Problema 5.5 (nella versione in avanti) ammette almeno una soluzione $\mathbf{y} \in C^1([t_0, t_0 + \tau]; \mathbb{R}^N)$.

Prova. Per $n \in \mathbb{N}$ suddivido l'intervallo $[t_0, t_0 + \tau]$ in n parti uguali introducendo la partizione

$$\mathcal{P}_n := \left\{ t_0, t_1 = t_0 + \frac{\tau}{n}, t_2 = t_0 + \frac{2\tau}{n}, \dots, t_n = t_0 + \tau \right\}.$$

Ovviamente, i punti t_i di $\mathcal{P}_n$ dipendono dalla scelta di n , ma per semplicità non evidenzieremo tale dipendenza nella notazione. Definiamo ora una funzione $\mathbf{y}_{(n)} : [t_0, t_0 + \tau] \rightarrow \mathbb{R}^N$ candidata ad approssimare la soluzione che stiamo cercando. Per cominciare, poniamo

$$\mathbf{y}_{(n)}(t) := \mathbf{y}_0 + \mathbf{f}(t_0, \mathbf{y}_0)(t - t_0) \quad \text{se } t \in [t_0, t_0 + \tau/n].$$

Poniamo, poi, $\mathbf{y}_1 := \mathbf{y}_0 + \tau \mathbf{f}(t_0, \mathbf{y}_0)/n$. Per costruzione, $\mathbf{y}_1$ è il limite sinistro della funzione lineare $\mathbf{y}_{(n)}$ per t tendente a t_1 . Iterando il procedimento, possiamo ora porre

$$\mathbf{y}_{(n)}(t) := \mathbf{y}_1 + \mathbf{f}(t_1, \mathbf{y}_1)(t - t_1) \quad \text{se } t \in [t_1, t_1 + \tau/n].$$

Definiamo quindi $\mathbf{y}_2$ come il limite per $t \rightarrow t_2^-$ di tale funzione e, via via, procedendo per induzione, in n passi estendiamo $\mathbf{y}_{(n)}$ fino a $t = t_n$ (dove le assegniamo il valore $\mathbf{y}_n$).

La funzione $\mathbf{y}_{(n)}$ costruita ha interessanti proprietà. Innanzitutto, è ovviamente lineare a tratti. Inoltre, essa è derivabile in $[t_0, t_0 + \tau] \setminus \mathcal{P}_n$ e la sua derivata vale

$$\mathbf{y}'_{(n)}(t) = \mathbf{f}(t_i, \mathbf{y}_i) \quad \text{se } t \in (t_i, t_{i+1}).$$

Battezzo $\phi_{(n)}$ tale derivata (nei punti di $\mathcal{P}_n$ ove essa a priori non è definita le assegno un valore arbitrario). Grazie al Lemma 5.10 la famiglia $\{\mathbf{y}_{(n)}\}$ è $(\sqrt{N}M)$ -equilipschitziana. Inoltre,

$$\mathbf{y}_{(n)}(t) = \mathbf{y}_0 + \int_{t_0}^t \phi_{(n)}(s) \, ds \quad \forall t \in [t_0, t_0 + \tau], \quad n \in \mathbb{N}. \quad (5.6)$$

Applicando Ascoli a $\{\mathbf{y}_{(n)}\}$, segue l'esistenza di una sottosuccessione, che per semplicità non ridenominiamo, uniformemente convergente a un limite $\mathbf{y}$.

Per mostrare che $\mathbf{y}$ è la soluzione, chiaramente basta far vedere che $\phi_{(n)}(\cdot)$ tende a $\mathbf{f}(\cdot, \mathbf{y}(\cdot))$ uniformemente in $[t_0, t_0 + \tau]$ e quindi passare al limite sotto integrale in

(5.6). Osserviamo allora che, per $t \in [t_i, t_{i+1})$,

$$\begin{aligned} |\mathbf{f}(t, \mathbf{y}(t)) - \phi_{(n)}(t)| &= |\mathbf{f}(t, \mathbf{y}(t)) - \mathbf{f}(t_i, \mathbf{y}_{(n)}(t_i))| \\ &\leq |\mathbf{f}(t, \mathbf{y}(t)) - \mathbf{f}(t, \mathbf{y}(t_i))| + |\mathbf{f}(t, \mathbf{y}(t_i)) - \mathbf{f}(t, \mathbf{y}_{(n)}(t_i))| \\ &\quad + |\mathbf{f}(t, \mathbf{y}_{(n)}(t_i)) - \mathbf{f}(t_i, \mathbf{y}_{(n)}(t_i))| \\ &=: I_1 + I_2 + I_3. \end{aligned} \tag{5.7}$$

Ora, notiamo che i grafici di $\mathbf{y}_{(n)}$ e $\mathbf{y}$, per costruzione, sono tutti contenuti nel “rettangolo” compatto $R := [t_0, t_0 + \tau] \times \overline{B}(\mathbf{y}_0, \sqrt{NM}\tau)$. Inoltre, tali funzioni sono $(\sqrt{NM})$ -equilipschitziane in $[t_0, t_0 + \tau]$. Poiché $\mathbf{f}$ è uniformemente continua su R grazie al teorema di Heine, fissando $\varepsilon > 0$ si osserva subito che, per n sufficientemente grande, $I_1 < \varepsilon$ grazie all’uniforme continuità di $\mathbf{y}$ (per n grande, t “si avvicina” a t_i); $I_2 < \varepsilon$ grazie alla convergenza uniforme di $\mathbf{y}_{(n)}$ a $\mathbf{y}$; infine, $I_3 < \varepsilon$ semplicemente per l’uniforme continuità di $\mathbf{f}$. Questo completa la prova del teorema. ■

Eliminiamo ora le ipotesi sulla limitatezza e sul dominio di $\mathbf{f}$:

Corollario 5.12. *Sia $\Omega \subset \mathbb{R}^{N+1}$ aperto, $(t_0, \mathbf{y}_0) \in \Omega$, $\mathbf{f} \in C(\Omega; \mathbb{R}^N)$. Allora il Problema di Cauchy (5.5) in avanti ammette almeno una soluzione in piccolo $(\mathbf{y}, T)$ ove $T > t_0$ è il tempo finale associato a $\mathbf{y}$ ai sensi della Def. 1.11.*

Prova. È sufficiente considerare una nuova funzione $\tilde{\mathbf{f}}$ che coincida con $\mathbf{f}$ su un “rettangolo” chiuso della forma $R := [t_0, t_0 + \tau] \times \overline{B}(\mathbf{y}_0, \varepsilon)$ e che sia definita, continua e limitata su $[t_0, t_0 + \tau] \times \mathbb{R}^N$. Si può costruire $\tilde{\mathbf{f}}$ prolungando in modo opportuno la restrizione di $\mathbf{f}$ a R . In questo modo, riscrivendo il Problema 5.5 con $\tilde{\mathbf{f}}$ al posto di $\mathbf{f}$, si può applicare il Teorema 5.11 che garantisce l’esistenza di almeno una soluzione $\tilde{\mathbf{y}}$. Dato che $\tilde{\mathbf{y}}(t_0) = \mathbf{y}_0$, per continuità $\tilde{\mathbf{y}}$, almeno per t sufficientemente vicino a t_0 (ossia fintanto che il suo grafico resta contenuto in R), deve risolvere anche il Problema 5.5 originario. ■

A questo punto, è naturale provare a ripercorrere la teoria svolta nel caso localmente Lipschitziano (Cap. 1). In particolare, si vede facilmente che possono essere estesi al caso senza unicità:

- la Def. 1.18 di prolungamento di una soluzione;
- la Def. 1.20 di soluzione massimale;
- la Prop. 1.22 (criterio di “non massimalità”);
- il Teorema 1.23 di esistenza e unicità “in grande”.

Ciò che invece cambia è il Teorema 1.21; infatti, l’insieme delle soluzioni del problema in avanti ora non sarà più *totalmente* ordinato rispetto alla relazione di estensione, ma varrà solo un ordinamento di tipo parziale. In quest’ottica, per dimostrare l’esistenza di una soluzione massimale non sarà più possibile procedere per costruzione diretta, ma occorrerà usare qualche strumento più sofisticato (ad esempio il seguente

Lemma 5.13. (Lemma di Zorn). *Sia $\mathcal{S}$ un insieme parzialmente ordinato, non vuoto e induttivo (ossia tale che ogni suo sottoinsieme totalmente ordinato ammette un maggiorante in $\mathcal{S}$). Allora $\mathcal{S}$ contiene almeno un elemento massimale.*

In particolare, si segnala che, data una soluzione in piccolo, questa in generale potrà essere prolungata da più soluzioni massimali distinte, che a loro volta potranno essere definite su domini diversi:

Esercizio 5.14. Studiare qualitativamente l'equazione

$$y' = \sqrt{|y|} + y^2.$$

5.4 Il “pennello di Peano”

Fino a questo momento, sappiamo che, anche in assenza di ipotesi di Lipschitz, il Problema (5.5) ammette almeno una soluzione. La mancanza di unicità è stata per ora evidenziata solo tramite la costruzione di esempi concreti. In questo paragrafo, vogliamo studiare più a fondo la struttura della famiglia delle soluzioni di (5.5), riferendoci per semplicità al caso scalare, e mostrando in particolare due risultati: il primo ci dirà che ogni soluzione è costruibile attraverso un procedimento di approssimazione; il secondo che esistono una soluzione “minima” $\underline{y}$ e una soluzione “massima” $\bar{y}$ tali che ogni altra soluzione y verifica $\underline{y} \leq y \leq \bar{y}$ almeno laddove tutte e tre le soluzioni sono definite. Inoltre, per ogni punto della regione di piano compresa tra i grafici di $\underline{y}$ e $\bar{y}$ passa almeno una soluzione. Questa è proprio una formalizzazione del fenomeno che si era chiamato “pennello di Peano”. Per semplicità, ci riferiremo sempre al problema in avanti e ci metteremo nelle ipotesi del Teorema 5.11, a cui peraltro ci si può sempre ricondurre, almeno localmente, tramite il procedimento evidenziato nella prova del Cor. 5.12. Veniamo dunque al primo risultato, che in realtà è un teorema per modo di dire, visto che l'enunciato di per sè è banalmente vero. Tuttavia, il procedimento “costruttivo” adoperato nella prova è di per sè interessante.

Teorema 5.15. *Sia, nelle ipotesi del Teorema 5.11, y una soluzione di (PCA). Allora esiste una successione $\{y_{(n)}\}$ di funzioni lineari a tratti che tende uniformemente a y .*

Prova. Osserviamo innanzitutto che y è M -Lipschitziana su $[t_0, t_0 + \tau)$, poiché $|f|$ è per ipotesi limitata da M e qui $N = 1$. Dunque, si ha che $\text{graf}(y) \subset [t_0, t_0 + \tau) \times [y_0 - M\tau, y_0 + M\tau) =: R$ (gli intervalli sono stati ora presi tutti semiaperti a destra per comodità nella suddivisione di R che faremo tra poco). Fissato $n \in \mathbb{N}$, suddividiamo R in $2n^2$ parti ponendo, per $i = 0, \dots, n-1$, $j = 0, \dots, 2n-1$,

$$R_{i,j} := [t_0 + i\tau/n, t_0 + (i+1)\tau/n) \times [y_0 - M\tau + jM\tau/n, y_0 - M\tau + (j+1)M\tau/n).$$

Ovviamente R è, per costruzione, l'unione degli $R_{i,j}$. A questo punto, dato $P \in R$, associamo a P un rettangolo $R(P)$ nel seguente modo:

- se P sta nell'interno di uno degli $R_{i,j}$ oppure sul suo bordo verticale (estremi di questo esclusi), poniamo $R(P) := R_{i,j}$;
- se P sta sul “bordo orizzontale” che divide $R_{i,j}$ da $R_{i,j+1}$ poniamo $R(P) := R_{i,j} \cup R_{i,j+1}$.

Poniamo ora:

$$M(P) := \max\{f(t, y) : (t, y) \in R(P)\}, \quad m(P) := \min\{f(t, y) : (t, y) \in R(P)\}.$$

A questo punto, sia $P_0 := (t_0, y_0)$ e sia μ_0 un *qualunque* numero nell'intervallo $[m(P_0), M(P_0)]$. Consideriamo il tratto di poligonale

$$y_{(n)}(t) := y_0 + \mu(P_0)(t - t_0)$$

definita per tutti i t tali per cui $y_n(t)$ “continua a stare” dentro $R(P_0)$. Se t_1 è l'istante in cui $y_{(n)}$ “tocca” la frontiera di $R(P_0)$, poniamo $y_1 := y_n(t_1)$ e $P_1 := (t_1, y_1)$. Iterando il procedimento, non è difficile mostrare che dopo un numero finito (detto $k + 1$) di passi (dipendente ovviamente dalla scelta che facciamo volta per volta delle pendenze μ_i , $i = 0, \dots, k$), otterremo che $t_k = t_0 + \tau$.

A questo punto, definendo $f_{(n)}(t, y)$ pari a $y'_{(n)}(t)$ nei punti del grafico di $y_{(n)}$ ove questa è derivabile e pari a f in tutto il resto di R , si vede subito che

$$y_{(n)}(t) = y_0 + \int_{t_0}^t f_{(n)}(s, y_{(n)}(s)) \, ds \quad \forall t \in [t_0, t_0 + \tau), \quad (5.8)$$

ossia $y_{(n)}$ risolve l'equazione di Volterra associata alla $f_{(n)}$ (si noti che $f_{(n)}$ non è in genere una funzione continua in R , ma è certamente integrabile; peraltro, la formula qui sopra “vede” $f_{(n)}$ solo nei punti del grafico di $y_{(n)}$).

Grazie allora all'uniforme continuità di f e al fatto che per $n \rightarrow \infty$ il diametro dei rettangoli $R_{i,j}$ tende a 0, si vede subito che $f_{(n)}$ tende a f uniformemente. Inoltre, dato che le $y_{(n)}$ sono M -equilipschitziane, per il teorema di Ascoli, una loro sottosuccessione tende uniformemente a un limite z . Passando al limite nella (5.8), si nota che tale funzione z è una soluzione di (5.5). Tuttavia, essa dipende dalla scelta che abbiamo fatto nei singoli passaggi delle “pendenze” μ_i . Quanto dunque ci resta da fare è mostrare che queste possono essere prese in modo tale che z sia proprio uguale alla specifica soluzione y che vogliamo approssimare.

Traccio dunque, all'interno del rettangolo R , il grafico di y ed evidenzio i punti in cui questo interseca le maglie della rete che genera la suddivisione negli $R_{i,j}$. Interpolando linearmente tali punti, si costruisce subito la poligonale “approssimante” $y_{(n)}$. Si noti che $y_{(n)}$ è una funzione lineare a tratti, in quanto, per la Lipschitzianità di y , i punti di intersezione di $\text{graf}(y)$ con le maglie sono “essenzialmente” in numero finito, nel senso che in realtà possono esservi infiniti punti “consecutivi” della stessa ordinata ma questi vengono ad essere “congiunti” da singoli tratti orizzontali dell'approssimante. Ora, utilizzando il Teorema di Lagrange, si può vedere che il grafico di $y_{(n)}$ ha in ogni punto (fatti salvi i “vertici”) pendenza compatibile con il procedimento generale di approssimazione delineato qui sopra. Più precisamente, se P_i e P_{i+1} sono due intersezioni consecutive, la pendenza di $y_{(n)}$ nel tratto che le congiunge è necessariamente un valore tra $m(P_i)$ e $M(P_i)$ e dunque può essere preso come μ_i . Grazie alla convergenza del procedimento di approssimazione, si ha dunque che $y_{(n)}$ converge proprio alla soluzione y scelta all'inizio. ■

Facciamo ora vedere che esiste una soluzione $\bar{y}$ maggiore di ogni altra soluzione. Ci serviamo del

Lemma 5.16. *Sia $I \subset \mathbb{R}$ un intervallo chiuso e limitato; $\mathcal{F} \subset C(I; \mathbb{R})$ una famiglia uniformemente limitata, uniformemente equicontinua e totalmente ordinata rispetto*

all'usuale ordinamento di funzioni. Posto allora, per $t \in I$,

$$\bar{\phi}(t) := \sup\{\phi(t) : \phi \in \mathcal{F}\},$$

si ha che $\forall \varepsilon > 0$ esiste $\phi_\varepsilon \in \mathcal{F}$ tale che $\|\bar{\phi} - \phi_\varepsilon\|_\infty \leq 5\varepsilon$.

Prova. Comincio a mostrare che $\bar{\phi}$ è anch'essa uniformemente continua. In corrispondenza di ε sia δ buono per l'uniforme equicontinuità di $\mathcal{F}$ e siano s, t due punti di I distanti meno di δ . È allora facile vedere che esiste $\phi_{s,t} \in \mathcal{F}$ tale che $\phi_{s,t}(s) \geq \bar{\phi}(s) - \varepsilon$ e $\phi_{s,t}(t) \geq \bar{\phi}(t) - \varepsilon$. Dunque,

$$|\bar{\phi}(s) - \bar{\phi}(t)| \leq |\bar{\phi}(s) - \phi_{s,t}(s)| + |\phi_{s,t}(s) - \phi_{s,t}(t)| + |\phi_{s,t}(t) - \bar{\phi}(t)| \leq 3\varepsilon.$$

Sia ora $\delta = \delta(\varepsilon)$ corrispondente alla condizione di uniforme equicontinuità. Siano $t_1, t_2, \dots, t_k \in I$ tali che $I \subset \cup_{i=1, \dots, k} (t_i - \delta, t_i + \delta)$. È facile allora costruire $\phi_\varepsilon \in \mathcal{F}$ tale che $\phi_\varepsilon(t_i) \geq \bar{\phi} - \varepsilon$ per ogni $i = 1, \dots, k$. Grazie all'uniforme continuità di $\bar{\phi}$ ed equicontinuità di $\mathcal{F}$, ϕ_ε è la funzione desiderata. ■

Consideriamo ora l'insieme $\mathcal{S}$ delle soluzioni di (5.5). Esso è parzialmente ordinato rispetto all'usuale relazione d'ordine di funzioni, equilimitato ed equilipschitziano. Dato allora un qualunque sottoinsieme totalmente ordinato $\mathcal{F} \subset \mathcal{S}$, il lemma precedente garantisce che $\mathcal{F}$ ammette un maggiorante $\bar{\phi} \in \mathcal{S}$ (si noti che il fatto che $\bar{\phi}$ sia anch'essa una soluzione si ottiene passando al limite nell'equazione di Volterra). Per il Lemma di Zorn, allora $\mathcal{S}$ ammette un elemento massimale, che chiamo $\bar{y}$. Inoltre, osservo che tale elemento massimale è unico. Se, infatti, $\bar{y}_1$ e $\bar{y}_2$ fossero elementi massimali distinti, posto $\bar{y}(t) := \max\{\bar{y}_1(t), \bar{y}_2(t)\}$ per $t \in [t_0, t_0 + \delta)$, avrei che $\bar{y}$ è ancora una soluzione (si noti che, in particolare, $\bar{y}$ è ancora una funzione di classe C^1), la quale sarebbe strettamente maggiore (nel senso dell'ordinamento di funzioni, ossia basta che lo sia in un punto!) sia di $\bar{y}_1$ che di $\bar{y}_2$, contraddicendo la massimalità di queste. Con lo stesso tipo di ragionamento, si mostra che $\bar{y}$ è confrontabile (e maggiore o uguale) rispetto a qualsiasi altro elemento di $\mathcal{S}$. Dunque, possiamo dire che $\bar{y}$ è la soluzione massima.

Notiamo anche che, grazie al Teorema 5.15, esiste una successione $y_{(n)}$ di poligonali convergente a $\bar{y}$; anzi, ti dovrebbe essere possibile dimostrare l'esistenza della soluzione massima costruendo opportunamente la $y_{(n)}$ e senza ricorrere al Lemma di Zorn.

Concludiamo con il

Teorema 5.17. (Pennello di Peano). *Siano $\bar{y}$ e $\underline{y}$ rispettivamente le soluzioni massima e minima del Problema in avanti (5.5) e sia P un qualunque punto appartenente alla regione strettamente compresa tra i loro grafici. Allora, P è contenuto nel grafico di almeno una soluzione di (5.5).*

Prova. Consideriamo il problema all'indietro con condizione iniziale fissata in P . Per il Teorema 5.11, tale problema ha almeno una soluzione in piccolo massimale. Tale soluzione o "incontra" una tra $\bar{y}$ e $\underline{y}$ (e, allora, si può incollare a questa raggiungendo così $P_0 = (t_0, y_0)$) o resta sempre compresa tra i loro grafici, e deve in tal caso comunque raggiungere P_0 (non può "fermarsi prima" perché altrimenti sarebbe violato il criterio di non massimalità). ■

Bibliografia

- [1] F. Conti, P. Acquistapace, A. Savojni, “Analisi Matematica – Teoria e Applicazioni”, McGraw-Hill Italia, 2001.
- [2] G. Gilardi, “Analisi uno”, seconda edizione, McGraw-Hill Italia, 1995.
- [3] G. Gilardi, “Analisi due”, seconda edizione, McGraw-Hill Italia, 1996.
- [4] G. Gilardi, “Analisi Matematica di Base”, McGraw-Hill Italia, 2001.
- [5] M.W. Hirsch, S. Smale, “Differential Equations, Dynamical Systems, and Linear Algebra”, Academic Press, New York, 1970.
- [6] S. Lang, “Algebra Lineare”, Boringhieri, 1970.
- [7] , L. Pistone, “Equazioni Differenziali Ordinarie in Ipotesi di Sola Continuità”, Tesi di Laurea, Università di Pavia, 2005.
- [8] , L.C. Piccinini, G. Stampacchia, G. Vidossich, “Equazioni Differenziali Ordinarie in $\mathbb{R}^N$ ”, Liguori, Napoli, 1978.
- [9] S. Salsa, A. Squellati, “Esercizi di Analisi Matematica 2 – Parte terza”, Masson, 1996.