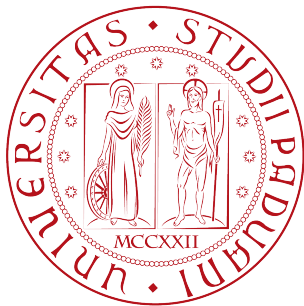




UNIVERSITÀ DEGLI STUDI DI PADOVA

Dipartimento di ingegneria dell'informazione

Corso di Laurea in INGEGNERIA ELETTRONICA

Elaborato di Laurea

Raggiungibilità dei sistemi dinamici e matrici di raggiungibilità in k passi

candidato:

Tommaso Di lauro

matricola **538832**

relatore:

Augusto Ferrante

Anno Accademico 2012/2013

Sommario

Il presente lavoro ha lo scopo di illustrare la struttura algebrica di un sistema dinamico e di investigare il problema della raggiungibilità.

In particolare ci si sofferma sui legami esistenti tra le variabili di ingresso e di uscita e le variabili di stato; l'obiettivo è dare una soluzione al problema di *determinare, a partire da uno stato iniziale, l'insieme degli stati in cui può essere guidato il sistema in corrispondenza degli ingressi ammissibili*. Si determineranno le condizioni necessarie e sufficienti per affermare quando uno stato è raggiungibile, grazie alla definizione di matrice di raggiungibilità

$$\mathcal{R} = [G \quad FG \quad F^2G \quad \dots \quad F^{k-1}G]$$

e del corrispondente teorema di Kalman; secondariamente si verificherà, grazie ad un recente teorema, *quando* una matrice nella suddetta forma è effettivamente una *matrice di raggiungibilità* per una coppia (F, G) di matrici.

Indice

1	Definizioni	2
1.1	Introduzione	2
1.2	Sistemi Dinamici	3
1.3	Definizione di stato	3
1.4	Mappa di aggiornamento dello stato	6
1.4.1	Forma implicita dei sistemi dinamici lineari	7
1.4.2	Forma esplicita delle evoluzioni	8
1.4.3	Equazioni di aggiornamento dello stato	8
1.5	Sistemi a eventi discreti	9
2	Raggiungibilità	11
2.1	Introduzione alla raggiungibilità	11
2.2	Raggiungibilità di sistemi dinamici lineari	12
2.2.1	Sistema dinamico a tempo continuo	12
2.2.2	Sistema dinamico a tempo discreto	14
2.3	Definizione di Raggiungibilità	15
2.3.1	Sistemi a tempo discreto	15
2.3.2	Sistemi a tempo continuo	16
3	Matrici di Raggiungibilità in k passi	18
3.1	definizioni	18
3.2	Matrici di raggiungibilità	18
A	Algebra lineare	21
A.1	polinomi	21
A.2	Proiezioni ortogonali	22
A.3	Matrici e sottospazi	24
A.4	Pseudoinversa di Moore-Penrose	25
A.4.1	definizioni	25
A.4.2	Relazioni tra i sottospazi	27
A.5	Esempi	28
A.5.1	Calcolo della Pseudoinversa	28
A.5.2	calcolo della matrice di proiezione	31
	Bibliografia	33

Capitolo 1

Definizioni

1.1 Introduzione

Un modello matematico è un insieme di teorie, di leggi e/o equazioni in grado di rappresentare (e descrivere) il funzionamento di un ente già presente nella realtà (oppure, ma è più raro, anche di un ente astratto, che deve cioè essere ancora costruito e/o osservato in natura).

I sistemi di interesse per questa trattazione non sono altro che strumenti matematici astratti (cioè modelli) che permettono di analizzare il comportamento di un sistema fisico e di controllare tale sistema secondo scopi predefiniti.

Il vero oggetto di interesse della seguente analisi, ovvero il cosiddetto sistema dinamico, sarà definito nel prossimo paragrafo.

Si restringerà da subito questa analisi su modelli nei quali la più importante (ancorché non unica) variabile indipendente sia il tempo, cioè nei quali sia sempre possibile individuare una qualche relazione tra le loro proprietà e il susseguirsi di istanti temporali distinti o distinguibili.¹ Questo è il motivo per il quale si porrà particolare interesse alla distinzione tra sistemi continui e discreti, nei quali la variabile tempo può appartenere a due insiemi numerici diversi. Si passerà poi alla definizione di stato (e conseguentemente a quella di variabile di stato) nella terza sezione e alla mappa di aggiornamento dello stato (cioè una sua utile rappresentazione matematica) nella quarta sezione.

Si osservi che i sistemi dinamici non sono l'unica alternativa possibile ai fini dello studio di un sistema fisico: nella quinta sezione verrà data una breve definizione dei sistemi ad eventi discreti (o DES). Nel capitolo secondo verrà introdotto il concetto di raggiungibilità per un sistema dinamico, con le opportune distinzioni tra sistemi continui e discreti. Nel capitolo terzo verrà presentato un recente risultato che permette di stabilire, mediante verifiche algebriche, quando un sistema è raggiungibile. In appendice sono riportati brevi richiami su concetti di algebra lineare, necessari a comprendere alcuni dettagli della trattazione.

¹Allo stato attuale di questa analisi urge stabilire che il tempo, in termini matematici, è un'entità misurabile. Esso è in un certo senso a sua volta un modello, cioè serve a stabilire una associazione tra il tempo di esperienza quotidiana e l'insieme di numeri, che si associano a istanti temporali. Entro un siffatto scenario si intende stabilire la possibilità di distinguere sempre un istante dal suo successivo o più in generale da un istante diverso.

1.2 Sistemi Dinamici

Una definizione puntuale, precisa e completa di sistema dinamico esula dagli scopi di questa indagine. Tuttavia per descrivere un sistema sono necessari almeno tre insiemi:

- $\mathbf{T}$, ovvero l'insieme dei tempi;
- $\mathbf{U}$, cioè un insieme delle variabili di ingresso;
- $\mathbf{Y}$, ovvero un insieme delle variabili di uscita;

L'insieme $\mathbf{T}$, cioè quello dei tempi può essere classificato in base alla sua cardinalità. Si configurano 2 casi di interesse:

- l'insieme è quello dei numeri reali $\mathbb{R}$, e in tale caso i modelli matematici utilizzabili si rifanno alle equazioni differenziali;
- l'insieme è quello dei numeri interi $\mathbb{Z}$ o in alternativa quello dei numeri naturali $\mathbb{N}$, e in questo caso i modelli matematici sono quelli delle equazioni alle differenze.

Nel primo caso il sistema sotto indagine prende il nome di sistema a **tempo continuo**; nel secondo invece il sistema è detto a **tempo discreto**.²

Gli insiemi $\mathbf{U}$ e $\mathbf{Y}$ normalmente sono funzioni del tempo, quindi l'insieme $\mathbf{U}$ degli ingressi assume la forma:

$$\mathbf{U} = \{u_1(t), u_2(t), u_3(t), \dots, u_n(t)\} \quad (1.1)$$

Gli ingressi sono entità che agiscono dall'esterno del sistema verso il sistema stesso, e che tipicamente determinano una sua evoluzione. L'insieme $\mathbf{Y}$ delle uscite sarà costituito da elementi del tipo

$$\mathbf{Y} = \{y_1(t), y_2(t), y_3(t), \dots, y_n(t)\} \quad (1.2)$$

Le uscite sono entità governate dal sistema che agiscono verso l'esterno rispetto al sistema stesso.

Per studiare un sistema dinamico è necessario individuare una relazione tra questi tre insiemi. Oltre agli insiemi appena descritti si utilizza un quarto insieme di variabili, che sono dette variabili di stato. Una definizione di stato viene data nel prossimo paragrafo.

1.3 Definizione di stato

Lo *stato* di un sistema ad un *preciso istante* t serve a descrivere la sua situazione in tale istante. Se ci si riferisce al modello definito come sistema dinamico, cioè una terna in cui sia presente almeno un insieme di ingressi $\mathbf{U}$, e un insieme di uscite $\mathbf{Y}$, il concetto di *stato* rappresenta un'informazione che riassume tutta l'evoluzione passata del sistema. Si vuole cioè che l'informazione aggiuntiva conferisca la possibilità di studiare l'evoluzione *futura* del sistema dinamico qualunque sia stato il suo *passato*. Ecco quindi la seguente

Definizione 1.3.1 (Stato di un sistema dinamico). Si consideri un sistema in cui $u(t)$ sono le variabili di ingresso e $y(t)$ le variabili di uscita. Il segnale $x(t)$ si dice stato del sistema se, $\forall t_0 \in \mathbf{T}$, la conoscenza di $x(t_0)$ e di $u(t) \forall t \geq t_0$ permette di determinare **sia** l'evoluzione di $y(t)$ **sia** lo stato $x(t)$ per $t \geq t_0$.

²Solitamente questa distinzione è di importanza rilevante. Ad esempio le moderne tecnologie che stanno alla base di un qualunque sistema digitale prevedono che il tempo sia quantizzato, cioè suddiviso in entità minime indivisibili. Conseguenza logica di ciò è che i sistemi di questo tipo saranno a tempo discreto. Come si vedrà questa analisi non potrà prescindere da questa classificazione.

In altre parole, per i sistemi sotto indagine vige la proprietà di **separazione delle variabili** di stato, e cioè che

Corollario 1.3.1. *Ai fini dello studio dell'evoluzione futura, l'evoluzione passata del sistema è irrilevante a patto di conoscere lo stato presente.*

Un modello efficace per lo studio dei sistemi dinamici è quello del cosiddetto **modello in forma di stato**, nel quale l'informazione più importante è proprio lo **stato** del sistema. Nel contesto di un simile modello è possibile individuare un legame tra l'insieme degli ingressi **U** (cioè lo spazio vettoriale che ha come base U) e l'insieme delle uscite **Y** (cioè lo spazio vettoriale che ha come basi **Y**).

Tale legame, nel caso di sistemi dinamici continui è esprimibile tramite il sistema

$$\begin{cases} \dot{\mathbf{x}} = f(x(t), u(t), t) \\ \mathbf{y}(t) = h(x(t), u(t), t) \end{cases} \quad (1.3)$$

nel quale $\mathbf{x}(t) \in \mathbb{R}^n$, e la notazione

$$\dot{\mathbf{x}}(t) = \frac{d\mathbf{x}(t)}{dt}$$

sta ad indicare l'operatore derivata eseguita su ogni componente del vettore $\mathbf{x}(t)$, mentre

$$\mathbf{y}(t) = \begin{bmatrix} y_1(t) \\ y_2(t) \\ \vdots \\ y_p(t) \end{bmatrix} \in \mathbb{R}^p, \forall t$$

sta ad indicare il vettore appartenente allo spazio delle uscite.

I vettori

$$\mathbf{x}(t) = \begin{bmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_n(t) \end{bmatrix}$$

sono detti **stati** del sistema, e lo spazio vettoriale $W \subset \mathbb{R}^n$ al quale essi appartengono è detto **spazio di stato**. Le componenti $x_1(t)$, $x_2(t)$, $x_n(t)$ di un particolare stato sono dette **variabili di stato**.

Si consideri che le variabili di stato sono ancora una volta dei modelli, cioè *rappresentano* in forma matematica alcune precise proprietà del sistema; tuttavia potrebbero non avere una corrispondenza diretta con entità fisicamente rilevabili sul sistema reale. Si può anche affermare che per uno stesso sistema non esiste un criterio per una scelta univoca delle variabili di stato. Tuttavia, in generale è preferibile associare ad esse una quantità fisica che tenga conto delle possibili variazioni energetiche che avvengono nel contesto del sistema.

Esempio 1.3.0.1 (serbatoio per liquidi). Si consideri il sistema individuato da un serbatoio riempito con un liquido. Si indichi con h il livello raggiunto da tale liquido sulle sue pareti, cioè l'altezza della colonna equivalente. Dalla legge di Stevino è ben noto che la pressione esercitata dal liquido sul fondo del contenitore è esprimibile con la formula

$$P_h = \rho \cdot g \cdot h \quad (1.4)$$

dove P_h è la pressione idrostatica esercitata dal liquido;

ρ è la densità del liquido;

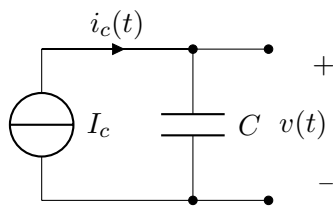
g è la costante di gravità, pari a $9.81 \frac{m}{s^2}$;

h è l'altezza della colonna di liquido.

La formula cioè permette di asserire che la pressione sul fondo del contenitore non dipende dal *volume totale* di liquido presente, ma piuttosto dall'altezza raggiunta dal liquido all'interno della struttura contenitiva. La pressione idrostatica esercitata è strettamente correlata all'energia potenziale del sistema.³ Una possibile scelta della variabile di stato potrebbe essere proprio l'altezza h , dal momento che tutte le altre proprietà del sistema fisico sono dei parametri da ritenersi costanti.

Un ulteriore semplice esempio può riferirsi ad un sistema elettrico.

Esempio 1.3.0.2 (condensatore elettrico). Si consideri un generatore di corrente impressa connesso in serie ad un condensatore, secondo lo schema di figura.



Il generatore di corrente I_c è un modello di un dispositivo in grado di imporre (nel ramo in cui esso è inserito) una corrente indipendente dal circuito circostante. Esso è connesso ad un elemento passivo in grado di accumulare carica elettrostatica, il condensatore elettrico: quest'ultimo è un modello di un elemento fisico di accumulo elettrostatico, il quale è in grado di instaurare tra le sue due armature una differenza di potenziale proporzionale alla separazione di carica elettrostatica subita. Il capacitore può fungere da serbatoio di carica, cioè può mantenere la carica accumulata anche dopo il distacco della sorgente. Denominata Q la carica accumulata, V la differenza di potenziale e C la capacità elettrica, questo elemento obbedisce alla seguente equazione

$$\frac{\Delta Q}{\Delta t} = C \frac{\Delta V}{\Delta t} \quad (1.5)$$

La quale, per $\Delta t \rightarrow 0$ diventa una equazione differenziale del primo ordine, del tipo

$$\frac{dq}{dt} = C \frac{dv}{dt} \quad (1.6)$$

dalla quale si ricava un legame diretto tra la corrente impressa e la differenza di potenziale

$$i(t) = C \frac{dv}{dt} \quad (1.7)$$

La differenza di potenziale è strettamente correlata all'energia accumulata nel sistema condensatore. Una buona scelta per le variabili di stato, quando l'elemento di accumulo è di questo tipo, è la differenza di potenziale $v(t)$.

³Più precisamente la pressione idrostatica, se espressa in [Pascal], è omogenea all'energia potenziale specifica posseduta da un elemento di volume infinitesimo ∂V posto ad un'altezza h . Infatti l'unità di misura $[Pa] = [\frac{N}{m^2}]$ mentre l'unità di misura dell'energia specifica, ovvero energia per unità di volume è $[\frac{J}{m^3}] = [\frac{N \cdot m}{m^3}] = [\frac{N}{m^2}]$.

1.4 Mappa di aggiornamento dello stato

Nel contesto di un sistema dinamico è opportuno definire la struttura delle funzioni definite tramite il sistema (1.3). Le funzioni f e h citate non hanno alcuna caratteristica peculiare. Tuttavia, i sistemi di maggiore interesse pratico sono quelli nei quali le funzioni godono di due specifiche proprietà.

Definizione 1.4.1. un sistema dinamico si dice tempo-invariante se le funzioni f e h non dipendono dal tempo.

Se un sistema possiede questa proprietà, nel caso di sistemi a tempo continuo si può scrivere la forma seguente:

$$\begin{cases} \dot{x}_1(t) = f_1(x_1(t), \dots, x_n(t), u_1(t), \dots, u_m(t)) \\ \vdots \\ \dot{x}_n(t) = f_n(x_1(t), \dots, x_n(t), u_1(t), \dots, u_m(t)) \end{cases} \quad (1.8)$$

e

$$\begin{cases} y(t) = h_1(x_1(t), \dots, x_n(t), u_1(t), \dots, u_m(t)) \\ \vdots \\ y_p(t) = h_p(x_1(t), \dots, x_n(t), u_1(t), \dots, u_m(t)) \end{cases} \quad (1.9)$$

Le prime assumono la denominazione di equazioni di stato; le seconde prendono il nome di equazioni di uscita.

Nel caso di sistemi a tempo discreto la struttura appare come la seguente:

$$\begin{cases} x_1(t+1) = f_1(x_1(t), \dots, x_n(t), u_1(t), \dots, u_m(t)) \\ \vdots \\ x_n(t+1) = f_n(x_1(t), \dots, x_n(t), u_1(t), \dots, u_m(t)) \end{cases} \quad (1.10)$$

e

$$\begin{cases} y(t) = h_1(x_1(t), \dots, x_n(t), u_1(t), \dots, u_m(t)) \\ \vdots \\ y_p(t) = h_p(x_1(t), \dots, x_n(t), u_1(t), \dots, u_m(t)) \end{cases} \quad (1.11)$$

Definizione 1.4.2. un sistema dinamico si dice lineare se le funzioni f e h sono funzioni lineari dei propri argomenti.

Questa affermazione si traduce nella possibilità di poter scrivere le equazioni di cui sopra come semplice combinazione lineare di addendi. Il vantaggio è enorme, e consente una rappresentazione di tipo algebrico.

1.4.1 Forma implicita dei sistemi dinamici lineari

Se un sistema possiede entrambe le proprietà è possibile riscrivere le equazioni nel modo seguente:

$$\begin{cases} \dot{x}_1(t) = f_{11}x_1(t) + \dots + f_{1n}x_n(t) + g_{11}u_1(t) + \dots + g_{1m}u_m(t) \\ \vdots \\ \dot{x}_n(t) = f_{n1}x_1(t) + \dots + f_{nn}x_n(t) + g_{n1}u_1(t) + \dots + g_{nm}u_m(t) \end{cases} \quad (1.12)$$

$$\begin{cases} y_1(t) = h_{11}x_1(t) + \dots + h_{1n}x_n(t) + d_{11}u_1(t) + \dots + d_{1m}u_m(t) \\ \vdots \\ y_p(t) = h_{p1}x_1(t) + \dots + h_{pn}x_n(t) + d_{p1}u_1(t) + \dots + d_{pm}u_m(t) \end{cases} \quad (1.13)$$

Mentre, nel caso di sistemi a tempo discreto:

$$\begin{cases} x_1(t+1) = f_{11}x_1(t) + \dots + f_{1n}x_n(t) + g_{11}u_1(t) + \dots + g_{1m}u_m(t) \\ \vdots \\ x_n(t+1) = f_{n1}x_1(t) + \dots + f_{nn}x_n(t) + g_{n1}u_1(t) + \dots + g_{nm}u_m(t) \end{cases} \quad (1.14)$$

$$\begin{cases} y_1(t) = h_{11}x_1(t) + \dots + h_{1n}x_n(t) + d_{11}u_1(t) + \dots + d_{1m}u_m(t) \\ \vdots \\ y_p(t) = h_{p1}x_1(t) + \dots + h_{pn}x_n(t) + d_{p1}u_1(t) + \dots + d_{pm}u_m(t) \end{cases} \quad (1.15)$$

Il sistema di equazioni relativo allo stato nel caso tempo-continuo in (1.12) può essere riscritto più brevemente come:

$$\dot{\mathbf{x}}(t) = F\mathbf{x}(t) + G\mathbf{u}(t) \quad (1.16)$$

oppure, nel caso tempo discreto in (1.14)

$$\mathbf{x}(t+1) = F\mathbf{x}(t) + G\mathbf{u}(t) \quad (1.17)$$

Cioè in un sistema di equazioni rispettivamente differenziali o alle differenze dove $\mathbf{x}(t)$, $\mathbf{u}(t)$, $\mathbf{y}(t)$ sono vettori di dimensione n , m e p rispettivamente; le matrici F , G , H e D sono del tipo

$$F = \begin{bmatrix} f_{11} & \dots & f_{1n} \\ f_{21} & \dots & f_{2n} \\ \vdots & & \\ f_{n1} & \dots & f_{nn} \end{bmatrix} \quad G = \begin{bmatrix} g_{11} & \dots & g_{1m} \\ g_{21} & \dots & g_{2m} \\ \vdots & & \\ g_{n1} & \dots & g_{nm} \end{bmatrix}$$

$$H = \begin{bmatrix} h_{11} & \dots & h_{1n} \\ h_{21} & \dots & h_{2n} \\ \vdots & & \\ h_{p1} & \dots & h_{pn} \end{bmatrix} \quad D = \begin{bmatrix} d_{11} & \dots & d_{1m} \\ d_{21} & \dots & d_{2m} \\ \vdots & & \\ d_{p1} & \dots & d_{pm} \end{bmatrix}$$

di dimensione $n \times n$, $n \times m$, $p \times n$ e $p \times m$.

Tali equazioni prendono il nome di *mappa di aggiornamento dello stato ad un passo* (nel caso di sistema a tempo discreto) o del **primo ordine** (nel caso di sistema a tempo continuo), ed è anche detta **forma implicita**.

1.4.2 Forma esplicita delle evoluzioni

Il legame ingresso-stato è qui rappresentato in forma implicita. Per ottenere la forma esplicita è opportuno risolvere le equazioni alle differenze (1.17) oppure le equazioni differenziali (1.16). Tali soluzioni hanno la forma

$$x(t) = x_l(t) + x_f(t) \quad (1.18)$$

Dove $x_l(t)$ è denominata evoluzione libera mentre $x_f(t)$ verrà denominata evoluzione forzata. Nel caso dei sistemi a tempo discreto le soluzioni possono essere esplicitate nel modo seguente:

$$x_l(k) \triangleq F^k x(0) \quad (1.19)$$

$$x_f(k) \triangleq \sum_{i=0}^{k-1} F^{k-1-i} G u(i) \quad (1.20)$$

Mentre nel caso dei sistemi a tempo continuo le soluzioni assumono la forma

$$x_l(t) = e^{Ft} x(0) \quad (1.21)$$

$$x_f(t) = \int_0^t e^{F(t-\sigma)} G u(\sigma) d\sigma \quad (1.22)$$

L'equazione di evoluzione libera permette di descrivere, in forma matematica, i valori assunti dalla (o dalle) variabili di stato ad un istante t fissato quando il sistema non è soggetto a ingressi. Viceversa l'equazione di evoluzione forzata consente di formulare una precisa relazione tra il vettore degli ingressi e quello degli stati.

1.4.3 Equazioni di aggiornamento dello stato

Consideriamo per semplicità la mappa di aggiornamento dello stato ad un passo (1.17):

$$\mathbf{x}(t+1) = F\mathbf{x}(t) + G\mathbf{u}(t)$$

Se si considera come istante iniziale l'istante $t = 0$ si possono scrivere le seguenti relazioni:

$$x(1) = Fx(0) + Gu(0) \quad (1.23a)$$

$$x(2) = Fx(1) + Gu(1) = F^2x(0) + FG u(0) + Gu(1) \quad (1.23b)$$

$$\vdots \quad (1.23c)$$

$$x(k) = F^k x(0) + F^{k-1} G u(0) + \dots + G u(k-1) \quad (1.23d)$$

L'ultima equazione rappresenta lo stato raggiunto al passo k -esimo, che può essere riscritto come

$$x(k) = \sum_{i=0}^{k-1} F^{k-1-i} G u(i) \quad (1.24)$$

che si può scrivere in forma matriciale

$$x(k) = \begin{bmatrix} G & FG & F^2G & \dots & F^{k-1}G \end{bmatrix} \begin{bmatrix} u(k-1) \\ u(k-2) \\ \vdots \\ u(0) \end{bmatrix} \quad (1.25)$$

1.5 Sistemi a eventi discreti

I sistemi dinamici fino a qui trattati hanno 2 caratteristiche fondamentali:

- sono sistemi a stato continuo. In altre parole le variabili di stato possono assumere valori in un sottoinsieme di $\mathbb{R}^n$ con la cardinalità del continuo.
- L'evoluzione di detti sistemi è governata dal tempo, ovvero le variabili che li caratterizzano sono in realtà funzioni o successioni nel tempo t .

Come prima affermato il modello dei sistemi dinamici non è l'unica soluzione per poter descrivere un sistema fisico. Si può considerare un modello alternativo con queste 2 nuove caratteristiche:

- lo **spazio di stato** è un insieme discreto.
- le transizioni da uno stato all'altro sono governate da eventi, e non dallo scorrere del tempo.

Una simile configurazione prende il nome di **Sistema a eventi discreti**, o **DES**. Per definire un DES sarà quindi sufficiente definire

1. un insieme degli stati. Esso sarà omogeneo a $\mathbb{N}$; tuttavia per semplicità d'ora in poi verrà preso in considerazione il caso in cui il numero di stati sia finito.

$$\mathbf{S} = \{ s_1, s_2, s_3, \dots, s_n \}$$

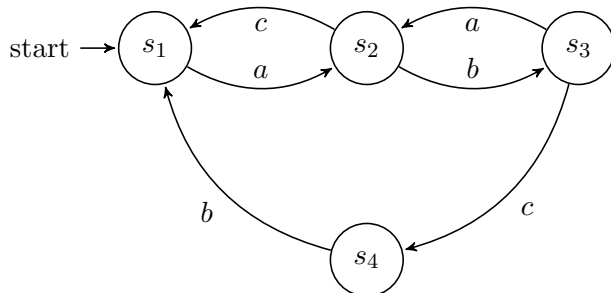
2. un cosiddetto insieme degli eventi o **spazio degli eventi**, ad esempio

$$\mathbf{E} = \{ a, b, c, \dots \}$$

vale a dire un insieme di transizioni possibili e

3. una **mappa delle transizioni**, cioè uno schema che associ ad ogni coppia di stati la o le transizioni che possono condurre da uno stato all'altro.

Una simile mappa si può facilmente rappresentare con un grafo orientato simile al seguente.



Si noti che l'insieme degli eventi potrebbe rappresentare semplicemente l'applicazione intenzionale di ingressi su un sistema dinamico del tipo definito precedentemente. Il modello DES non prevede l'utilizzo di equazioni differenziali: infatti qualunque fenomeno temporale è irrilevante ai fini del susseguirsi degli eventi, che divengono invece le uniche cause di evoluzione.

Tale modello è utilizzato per trattare i cosiddetti sistemi a stati finiti.

Definizione 1.5.1 (Sistema ad eventi discreti). Un sistema ad eventi discreti si definisce come una quintupla $M = (I, U, S, f, g)$, dove

$I = \{i_1, i_2, \dots, i_n\}$ insieme finito dei possibili simboli in ingresso

$U = \{u_1, u_2, \dots, u_m\}$ insieme finito dei possibili simboli in uscita

$S = \{s_1, s_2, \dots, s_h\}$ insieme finito degli stati

$f : I \times S \rightarrow S$ funzione di transizione (degli stati interni successivi); tale funzione mette in relazione lo stato nell'istante successivo al valore attuale dell'ingresso e dello stato, secondo l'equazione

$$S(t+1) = f(S(t), I(t))$$

$g : I \times S \rightarrow U$ funzione delle uscite (eventualmente parziale), che mette in relazione l'uscita al valore attuale dell'ingresso e dello stato, secondo l'equazione

$$U(t) = g(I(t), S(t))$$

Capitolo 2

Raggiungibilità

2.1 Introduzione alla raggiungibilità

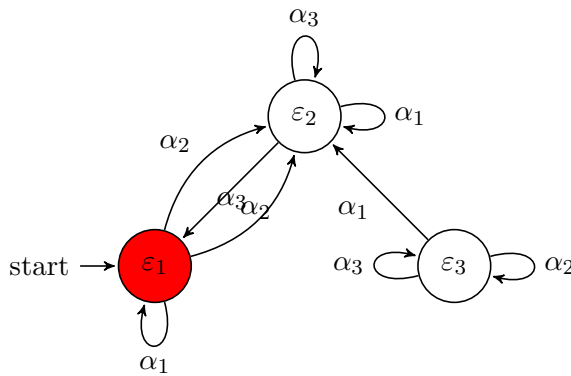
Sia dato un sistema con insieme degli stati non vuoto. Se si ammette esista uno stato iniziale, è ammissibile il problema di determinare l'insieme degli stati in cui può essere portato il sistema in corrispondenza degli ingressi (ovvero degli elementi appartenenti all'insieme degli ingressi), compatibili con quello stato. Un problema di questo tipo è detto *raggiungibilità*. Per converso, se si considera un stato finale desiderato il problema della *controllabilità* consiste nel determinare l'insieme di stati a partire dai quali esiste un ingresso che consenta di guidare il sistema nello stato voluto. Si consideri, per una prima analisi, un sistema a stati finiti ed eventi discreti. Sia $\mathbf{A}$ detto sistema a stati finiti. Il suo insieme S , o insieme degli stati sia

$$\mathbf{S} = \{ \varepsilon_1, \varepsilon_2, \varepsilon_3 \}$$

Il suo insieme E , o spazio degli eventi sia

$$\mathbf{E} = \{ \alpha_1, \alpha_2, \alpha_3 \}$$

Detto sistema ha una mappa delle transizioni¹ rappresentato mediante il grafo in figura e stato iniziale ε_1 , segnalato in rosso.



Un sistema così semplice si presta ad una analisi qualitativa.

Ispezionando il grafo si osserva che l'insieme degli stati raggiungibili dallo stato ε_1 , cioè dallo stato iniziale è

$$\Phi_1 = \{ \varepsilon_1, \varepsilon_2 \}$$

¹Si noti la presenza di autoanelli, cioè transizioni che portano il sistema verso lo stesso stato dal quale si era partiti.

mentre quello degli stati raggiungibili da ε_2 è

$$\Phi_2 = \{\varepsilon_1, \varepsilon_2\}$$

Invece, quello degli stati raggiungibili da ε_3 è

$$\Phi_3 = \{\varepsilon_3, \varepsilon_2\}$$

Con lo stesso principio è possibile asserire che l'insieme degli stati controllabili in un passo allo stato ε_1 è

$$\Gamma_1 = \{\varepsilon_1, \varepsilon_2\}$$

Mentre, per esempio, l'insieme degli stati controllabili in un passo dallo stato ε_2 è

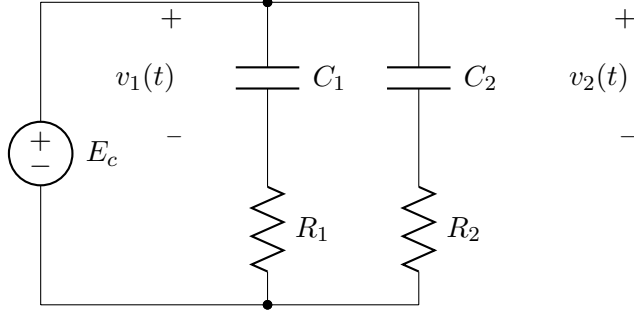
$$\Gamma_2 = \{\varepsilon_1, \varepsilon_2, \varepsilon_3\}$$

Come si vede dall'esempio, in generale gli insiemi di controllabilità e di raggiungibilità di uno stesso stato possono non coincidere.

2.2 Raggiungibilità di sistemi dinamici lineari

2.2.1 Sistema dinamico a tempo continuo

Esempio 2.2.1.1 (sistema elettrico). A titolo di esempio, si consideri il sistema dinamico ad un solo ingresso e due uscite corrispondente al circuito di figura. Si tratta di due capacità che vengono caricate mediante resistenze dalla stessa sorgente di tensione costante.



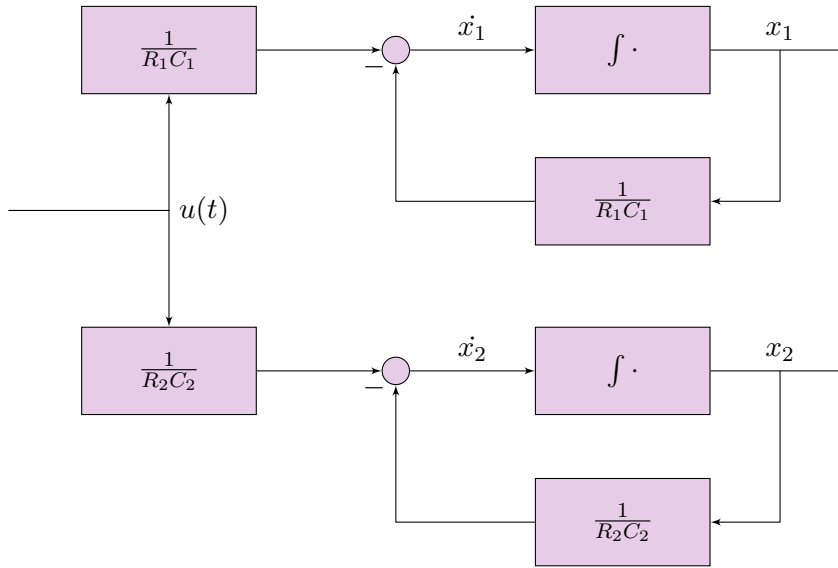
Per semplicità si considereranno ideali i componenti fisici connessi al circuito, cioè si supporranno concentrati in uno spazio limitato i parametri C , cioè la capacità e R , la resistenza elettrica e si assumerà che E_c sia un generatore ideale di tensione. Scegliendo come variabili di stato le tensioni ai capi dei condensatori, il vettore di stato del sistema è

$$\begin{bmatrix} x_1(t) \\ x_2(t) \end{bmatrix} = \begin{bmatrix} v_1(t) \\ v_2(t) \end{bmatrix} \quad (2.1)$$

Che per semplicità si considererà come vettore delle variabili di uscita. Le equazioni di aggiornamento dello stato sono una coppia di equazioni differenziali del primo ordine, che possono essere riassunte in forma di matrice:

$$\begin{bmatrix} \dot{x}_1(t) \\ \dot{x}_2(t) \end{bmatrix} = \begin{bmatrix} -\frac{1}{R_1 C_1} & 0 \\ 0 & -\frac{1}{R_2 C_2} \end{bmatrix} \begin{bmatrix} x_1(t) \\ x_2(t) \end{bmatrix} + \begin{bmatrix} \frac{1}{R_1 C_1} \\ \frac{1}{R_2 C_2} \end{bmatrix} u(t) \quad (2.2)$$

Si noti che è possibile ricondurre lo schema elettrico nel seguente schema a blocchi:



È evidente che i possibili stati raggiungibili dal sistema corrispondente alla rete di figura sono solo un sotto-insieme degli stati possibili. Per esempio il vettore (2.1) non potrà mai raggiungere lo stato

$$\mathbf{x}(t) = \begin{bmatrix} 0 \\ E \end{bmatrix} \quad (2.3)$$

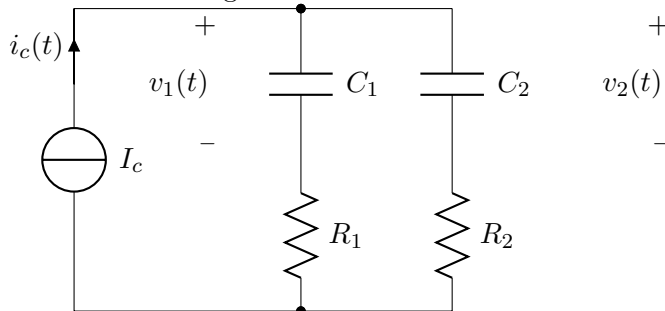
oppure, equivalentemente, lo stato

$$\mathbf{x}(t) = \begin{bmatrix} E \\ 0 \end{bmatrix} \quad (2.4)$$

Anche intuitivamente la conoscenza della teoria delle reti elettriche permetterebbe di stabilire che le due capacità sono destinate (per $t \rightarrow \infty$) ad essere caricate alla stessa tensione E_c , e ciò esclude che una delle due rimanga a tensione nulla. Un sistema così configurato presenta cioè dei limiti che dipendono dalla topologia della rete. Nel caso in esame è evidente che l'impossibilità di raggiungere gli stati (2.3) e (2.4) è legata al fatto che l'ingresso (cioè la sorgente di tensione costante) è applicato ad entrambi gli elementi di accumulo secondo lo stesso paradigma, ovvero secondo la schema parallelo.

Si noti che gli stati (2.3) e (2.4) appartengono ai sottospazi generati da $\begin{pmatrix} 1 \\ 0 \end{pmatrix}$ e $\begin{pmatrix} 0 \\ 1 \end{pmatrix}$.

Esempio 2.2.1.2 (sistema elettrico). Si consideri il caso in cui il generatore di tensione sia sostituito con un generatore di corrente. Il circuito diviene:



Anche in questo caso si assuma che le variabili di stato siano assegnate alle tensioni sulle capacità C_1 e C_2 . Considerando le equazioni di stato, si ottiene la coppia:

$$\begin{cases} \dot{x}_1(t) &= \frac{x_2(t)-x_1(t)}{(R_1+R_2)C_1} + u(t) \cdot \frac{R_2}{(R_1+R_2)C_1} \\ \dot{x}_2(t) &= \frac{x_1(t)-x_2(t)}{(R_1+R_2)C_2} + u(t) \cdot \frac{R_1}{(R_1+R_2)C_2} \end{cases}$$

Che conducono alla seguente matrice:

$$\begin{bmatrix} \dot{x}_1 \\ \dot{x}_2 \end{bmatrix} = \begin{bmatrix} -\frac{1}{C_1(R_1+R_2)} & \frac{1}{C_1(R_1+R_2)} \\ \frac{1}{C_2(R_1+R_2)} & -\frac{1}{C_2(R_1+R_2)} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} \frac{R_2}{(R_1+R_2)C_1} \\ \frac{R_1}{(R_1+R_2)C_2} \end{bmatrix} u$$

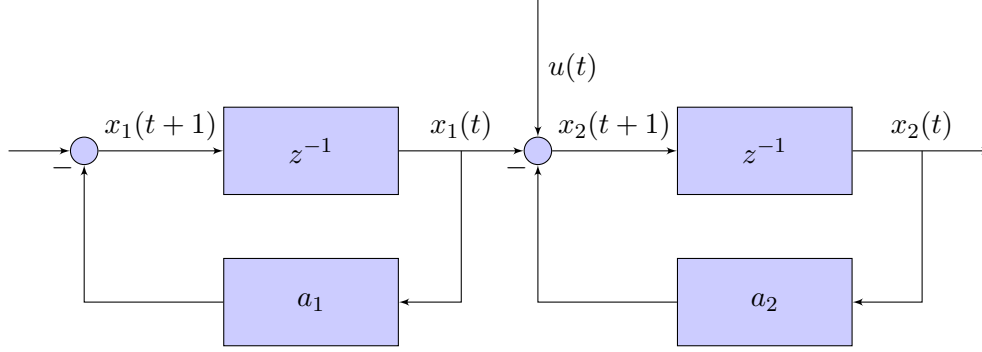
Fissate le costanti

$$a_1 = \frac{1}{C_1(R_1 + R_2)} \quad a_2 = \frac{1}{C_2(R_1 + R_2)}$$

si noti che le colonne della matrice quadrata non sono linearmente indipendenti: esse sono generate dallo stesso sottospazio vettoriale, che è $\begin{pmatrix} a_1 \\ -a_2 \end{pmatrix}$.

2.2.2 Sistema dinamico a tempo discreto

Si consideri, a titolo d'esempio, il sistema dinamico discreto di figura.



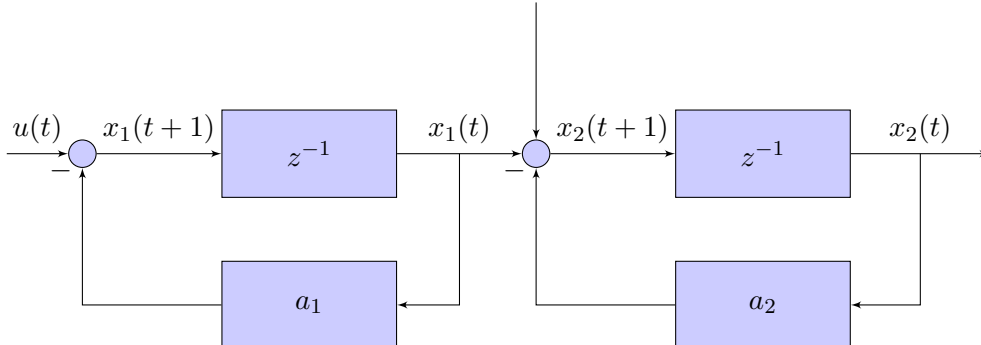
Le cui equazioni di stato sono

$$\begin{aligned} x_1(t+1) &= a_1 x_1(t) \\ x_2(t+1) &= x_1(t) + a_2 x_2(t) + u(t) \end{aligned}$$

Tali espressioni possono venire tradotte in forma di matrice

$$\begin{bmatrix} x_1(t+1) \\ x_2(t+1) \end{bmatrix} = \begin{bmatrix} a_1 & 0 \\ 1 & a_2 \end{bmatrix} \cdot \begin{bmatrix} x_1(t) \\ x_2(t) \end{bmatrix} + \begin{bmatrix} 0 \\ 1 \end{bmatrix} u(t) \quad (2.5)$$

Se $u(t)$ agisse all'ingresso del primo sistema, secondo lo schema seguente:



le equazioni diverrebbero

$$\begin{bmatrix} x_1(t+1) \\ x_2(t+1) \end{bmatrix} = \begin{bmatrix} a_1 & 0 \\ 1 & a_2 \end{bmatrix} \cdot \begin{bmatrix} x_1(t) \\ x_2(t) \end{bmatrix} + \begin{bmatrix} 1 \\ 0 \end{bmatrix} u(t) \quad (2.6)$$

È evidente che mentre nel primo caso gli stati raggiungibili sono il sotto-spazio vettoriale generato da $\langle (0, 1) \rangle$, nel secondo tutti gli stati sono raggiungibili al più in 2 passi. Si intuisce insomma che la effettiva raggiungibilità di uno stato dipenda sia dalla dinamica di aggiornamento degli stati (cioè dalla matrice F) sia dal modo in cui agiscono gli ingressi (cioè dalla matrice G).

Il problema su cui si vorrebbe indagare è se sia possibile definire condizioni necessarie e sufficienti per descrivere la raggiungibilità di un *qualunque* sistema, a prescindere da considerazioni di carattere topologico e strutturale; si vuole insomma astrarre rispetto alla realtà fisica o alla natura da cui deriva il sistema, considerando solo l'insieme di equazioni che lo descrivono.

2.3 Definizione di Raggiungibilità

2.3.1 Sistemi a tempo discreto

Per la classe dei sistemi a tempo discreto si possono stabilire condizioni di natura algebrica sullo spazio di stato.

All'interno dell'equazione (1.25) si può isolare la matrice

$$\mathcal{R}_k = [G \quad FG \quad F^2G \quad \dots \quad F^{k-1}G] \quad (2.7)$$

Detta matrice test di Kalman, che appartiene a $\mathbb{R}^{n \times nm}$. La sua immagine, ovvero

$$x_k^{\mathcal{R}} = \text{Im} [G \quad FG \quad F^2G \quad \dots \quad F^{k-1}G] \quad (2.8)$$

ha la struttura di uno spazio vettoriale. Esso è un sottospazio di $\mathbb{R}^n$, ed è l'insieme degli stati raggiungibili in k passi.

Definizione 2.3.1. il sottospazio $\in \mathbb{R}^n$ definito da $x_k^R \triangleq \text{Im} [G \quad FG \quad F^2G \quad \dots \quad F^{k-1}G]$ è detto sottospazio raggiungibile.

Si può osservare che i sotto-spazi raggiungibili in 1, 2, k passi soddisfano la catena

$$x_1^R \subseteq x_2^R \dots \subseteq x_k^R \quad (2.9)$$

Tale catena può divenire stazionaria, cioè del tipo

$$x_n^R = x_{n+1}^R \dots \triangleq x^R \quad (2.10)$$

se il numero $k > n$ ². In tale caso ogni stato è raggiungibile in un numero $k \leq n$ passi. A questo punto si può definire la raggiungibilità osservando che

Osservazione 1. Se $x^R = x$ allora il sistema è completamente raggiungibile.

E si può da subito formulare la seguente

Definizione 2.3.2. un sistema si dice raggiungibile (o completamente raggiungibile) se e solo se

$$\text{rank}(\mathcal{R}) = n$$

Cioè se il rango per colonne della matrice di raggiungibilità $\mathcal{R}$ è pieno, cioè ha le stesse dimensioni della matrice stessa.

²tale risultato è verificabile grazie al teorema di Cayley Hamilton, trattato in appendice

2.3.2 Sistemi a tempo continuo

Per quanto riguarda i sistemi continui è opportuno osservare che lo spazio vettoriale delle funzioni di ingresso non ha dimensione finita. Si può cioè notare che i valori assumibili da un qualsivoglia ingresso non sono successioni ($\in \mathbb{N}$) bensì funzioni ($\in \mathbb{R}$). Diviene indispensabile utilizzare l'equazione (1.21). Si introduce ora un apposito operatore lineare, definito come:

$$R_t : u(\cdot) \mapsto \int_0^t e^{F(t-\sigma)} G u(\sigma) d\sigma \quad (2.11)$$

e un apposito operatore aggiunto, definito di seguito:

$$R_t^* : X \rightarrow \mathcal{U} : x \mapsto u(\cdot) \quad (2.12)$$

Ed è finalmente lecito formulare la seguente

Definizione 2.3.3. Un generico stato x è raggiungibile se e solo se appartiene all'immagine di R_t , cioè al sottospazio raggiungibile al tempo t .

x è raggiungibile $\iff x \in \text{Im}(R_t)$. Il cosiddetto sottospazio raggiungibile, nei sistemi continui non dipende dalla lunghezza dell'intervallo $[0, t)$. Tale affermazione è giustificata dal prossimo

Teorema 2.3.1. Il sottospazio raggiungibile al tempo $t > 0$ è l'immagine della matrice di raggiungibilità, cioè di

$$\mathcal{R} = [G \quad FG \quad F^2G \quad \dots \quad F^{k-1}G]$$

Dimostrazione 2.3.1. Per le proprietà definite dall'operatore aggiunto si ha

$$\text{Im} R_t = (\ker R_t^*)^\perp \quad (2.13)$$

e

$$\text{Im}(\mathcal{R}) = (\ker \mathcal{R}^\top)^\perp \quad (2.14)$$

Per dimostrare che il sottospazio raggiungibile coincide con $\text{Im}(\mathcal{R})$ bisogna verificare che $\ker R_t^* = \ker \mathcal{R}^\top$. ciò è valido se:

$$\begin{aligned} x &\in \ker R_t^* \\ 0 &= G^\top e^{F^\top(t-\sigma)} x \\ 0 &= \sum_{i=0}^{\infty} G^\top (F^\top)^i x \frac{(t-\sigma)^i}{i!} \end{aligned} \quad (2.15)$$

$$0 = G^\top (F^\top)^i x \quad i = 0, 1, 2 \quad (2.16)$$

$$0 = G^\top (F^\top)^i x \quad i = 0, 1, 2, \dots, n-1 \quad (2.17)$$

$$x \in \ker \begin{bmatrix} G^\top \\ G^\top (F^\top) \\ \vdots \\ G^\top (F^\top)^{n-1} \end{bmatrix} = \ker(\mathcal{R})^\top$$

Osservazione 2. la condizione di raggiungibilità per uno stato di un sistema, che è data da

$$\text{Im}(\mathcal{R}) = x$$

corrisponde alla condizione

$$\ker \mathcal{R}^\top = \{0\}$$

Sia nei sistemi continui sia in quelli discreti il sottospazio X^R è l'immagine della matrice di raggiungibilità; si conclude che tutte le proprietà che si possono desumere solo dalla struttura di $\mathcal{R}$ valgono allo stesso modo per ambo le tipologie di sistema.

Esempio 2.3.2.1. Dato il sistema a tempo discreto descritto dalle matrici

$$x(t+1) = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{bmatrix} x(t) + \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} u(t)$$

si determinino i sotto-spazi di raggiungibilità in 3 passi.

Soluzione.

Si osservi che il sistema può dirsi raggiungibile perchè la matrice

$$F = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{bmatrix}$$

ha rango³ pieno. A partire dalla matrice test di Kalman si osserva che

$$\begin{aligned} x_1^R &= \text{Im}[G] = \text{span} \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} = \left\langle \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \right\rangle \\ x_2^R &= \text{Im} \begin{bmatrix} G & FG \end{bmatrix} = \left\langle \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right\rangle \\ x_3^R &= \text{Im} \begin{bmatrix} G & FG & F^2G \end{bmatrix} = \left\langle \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \right\rangle \end{aligned}$$

³ Per matrici non quadrate il *rango* di una matrice è $\text{rk}(A) = \dim \mathcal{R}(A) = \dim \mathcal{R}(A^T)$ cioè il rango è il minimo tra m, n . Le dimensioni degli spazi vettoriali sono legate dal teorema del rango e della nullità: per ogni matrice $A \in \mathbb{R}^{m \times n}$:

$$\dim(\text{Im}(A)) + \dim(\ker(A)) = n$$

oppure, equivalentemente:

$$\text{rk}(A) + \text{null}A = n$$

Capitolo 3

Matrici di Raggiungibilità in k passi

3.1 definizioni

Si definiscano le matrici $A \in K^{q \times q}$ e $B \in K^{q \times m}$ con $\mathbb{K}$ un campo algebrico. La matrice

$$R^{(k)}(A, B) = [B \quad AB \quad A^2B \quad \dots \quad A^{k-1}B] \in \mathbb{K}^{q \times k \cdot m} \quad (3.1)$$

è detta *matrice di raggiungibilità in k passi* della coppia (A, B) .

3.2 Matrici di raggiungibilità

In tempi recenti è stato formulato un risultato che permette di stabilire *quando* una matrice nella forma (3.1) è effettivamente una matrice di raggiungibilità per una coppia (A, B) di matrici. Tale risultato si può tradurre nel seguente

Teorema 3.2.1. *Sia data*

$$\mathcal{R} := [R_0 \quad R_1 \quad R_2 \quad \dots \quad R_{k-1}] \in \mathbb{R}^{n \times m \cdot k} \quad (3.2)$$

con $R_i \in \mathbb{R}^{n \times m}$. Si definiscano

$$\mathcal{R}_- := [R_0 \quad R_1 \quad R_2 \quad \dots \quad R_{k-2}] \in \mathbb{R}^{n \times m \cdot (k-1)}$$

e

$$\mathcal{R}^- := [R_1 \quad R_2 \quad R_3 \quad \dots \quad R_{k-1}] \in \mathbb{R}^{n \times m \cdot (k-1)}$$

Allora $\mathcal{R}$ è una matrice di raggiungibilità in k passi della coppia (A, B) , con $A \in \mathbb{R}^{n \times n}$ e $B \in \mathbb{R}^{n \times m}$ se e solo se

$$\ker[\mathcal{R}_-] \subseteq \ker[\mathcal{R}^-] \quad (3.3)$$

In questo caso la matrice $B = R_0$ e tutte le matrici A possono essere scritte nella forma¹

$$A = \mathcal{R}^- \mathcal{R}_-^\top (\mathcal{R}_- \mathcal{R}_-^\top)^\# + V \left[I - \mathcal{R}_- \mathcal{R}_-^\top (\mathcal{R}_- \mathcal{R}_-^\top)^\# \right] \quad (3.4)$$

Dove il parametro $V \in \mathbb{R}^{n \times n}$, $I = \begin{bmatrix} 1 & 0 & \dots \\ 0 & 1 & \dots \\ \vdots & \vdots & \ddots \end{bmatrix}$ è la matrice identità $I \in \mathbb{R}^{n \times n}$,

mentre $\mathcal{R}^\top$ indica la trasposta di $\mathcal{R}$.

¹La notazione $(\cdot)^\#$ indica la pseudo-inversa di Moore-Penrose, trattata in appendice.

Dimostrazione 3.2.1. Per ottenere il risultato bisogna dimostrare entrambe le implicazioni. Si osservi innanzi tutto che i prodotti

$$D = \mathcal{R}_- \mathcal{R}_-^\top \quad (3.5a)$$

$$E = \mathcal{R}^- \mathcal{R}^{-\top} \quad (3.5b)$$

devono dare luogo a matrici quadrate: infatti $\mathcal{R}_-$ e $\mathcal{R}^-$ hanno le medesime dimensioni, in modo tale che $D, E \in \mathbb{R}^{n \times n}$ cioè hanno le stesse dimensioni di A . Inoltre la matrice D è simmetrica, perchè risulta ottenuta dal prodotto di una matrice per la sua trasposta. Con la notazione appena introdotta si può riscrivere l'espressione (3.4) come segue:

$$A = ED^\# + V \left[I - DD^\# \right] \quad (3.6)$$

Per semplicità di consideri dapprima il caso particolare in cui le matrici D ed E siano invertibili, di modo che la pseudo inversa $D^\#$ coincida con la inversa, ovvero D^{-1} . l'equazione (3.6) si può riscrivere come:

$$A = ED^{-1} + V \left[I - DD^{-1} \right]$$

Grazie alle proprietà delle matrici inverse si ha, $\forall$ matrice $M \in \mathbb{R}^{n \times n}$ l'uguaglianza

$$MM^{-1} = M^{-1}M = I$$

e quindi gli elementi racchiusi entro parentesi quadre danno luogo alla matrice nulla. L'equazione si può nuovamente riscrivere come

$$A = ED^{-1}$$

Si osservi che questa equazione, detta *formula di Cramer*, è la soluzione al problema

$$\mathbf{A}\mathbf{D} = \mathbf{E}$$

in cui A sia una matrice incognita.²In generale la cancellazione della porzione tra parentesi avviene se D non è singolare.

Si supponga dunque che $\mathcal{R}$ sia effettivamente la matrice di raggiungibilità in k -passi della coppia di matrici A, B . Si potrà affermare che

$$A\mathcal{R}_- = \mathcal{R}^- \quad (3.7)$$

$$A \begin{bmatrix} R_0 & R_1 & R_2 & \dots & R_{k-2} \end{bmatrix} = \begin{bmatrix} R_1 & R_2 & \dots & R_{k-1} \end{bmatrix} \quad (3.8)$$

Ciascuna matrice R_k verrà premoltiplicata per A ; è evidente che in questo modo è assicurata l'inclusione (3.3). Per ottenere la (3.4) sarà sufficiente scrivere il problema

$$(A\mathcal{R}_-)^\top = (\mathcal{R}^-)^\top$$

²Il problema converso, cioè la scrittura

$$\mathbf{D}\mathbf{A} = \mathbf{E}$$

con A una matrice incognita, ha come soluzione l'espressione

$$A = D^{-1}E$$

E risolvere rispetto ad $A^\top$.

Per converso, si assuma che $\ker[\mathcal{R}_-] \subseteq \ker[\mathcal{R}^-]$.

Si osservi che l'*operatore di proiezione*³ sul sottospazio $\mathcal{R}_-^\top$ è definito dalla matrice

$$P_- \triangleq \mathcal{R}_-^\top \left(\mathcal{R}_- \mathcal{R}_-^\top \right)^\# \mathcal{R}_- \quad (3.9)$$

Tra le proprietà dell'operatore di proiezione $\mathbf{P}$ di una matrice A ve ne è una che permette di affermare che

$$\mathbf{P} \text{ proietta su } \text{Im}(A) \implies \text{Im}(\mathbf{P}) = \text{Im}(A) \quad (3.10)$$

Grazie alla (3.3) e tenendo conto di questo ultimo fatto si può riscrivere l'inclusione come:

$$P_- = \text{Im}(\mathcal{R}_-^\top) = (\ker(\mathcal{R}_-))^\perp \supseteq (\ker(\mathcal{R}^-))^\perp = \text{Im}((\mathcal{R}^-)^\top) \quad (3.11)$$

Si consideri ora la formula (3.4). Trascurando per il momento la porzione racchiusa tra parentesi quadrate, si può scrivere che

$$A = \mathcal{R}^- \mathcal{R}_-^\top \left(\mathcal{R}_- \mathcal{R}_-^\top \right)^\#$$

Allora vale la seguente catena di uguaglianze:

$$A \mathcal{R}_- = \mathcal{R}^- \underbrace{\mathcal{R}_-^\top \left(\mathcal{R}_- \mathcal{R}_-^\top \right)^\# \mathcal{R}_-}_{=P_-} = \mathcal{R}^- P_- \quad (3.12)$$

Se è verificata l'inclusione (3.11) l'esecuzione della proiezione ortogonale sul sotto-spazio $\mathcal{R}^-$ porta allo stesso sottospazio, quindi si può scrivere:

$$\mathcal{R}^- P_- = \mathcal{R}^- \quad (3.13)$$

confrontando le ultime due equazioni si deduce la (3.7).

Per accertare che la porzione racchiusa tra parentesi quadrate sia nulla, sarà sufficiente applicare la (3.7) sulla (3.4) (si trascuri la prima porzione della formula):

$$V \left[I - \mathcal{R}_- \mathcal{R}_-^\top \left(\mathcal{R}_- \mathcal{R}_-^\top \right)^\# \right] \mathcal{R}_- = V[\mathcal{R}_- - \mathcal{R}_- P_-] = V[\mathcal{R}_- - \mathcal{R}_-] = 0 \quad (3.14)$$

Si osservi che la porzione scritta tra parentesi quadre, ovvero $[I - DD^\#]$ non è altro che $(I - P)$, ovvero il *proiettore ortogonale* nel nucleo della trasposta di D , cioè in $\ker(D^\top)$.

³Tattato in appendice. Si veda la (A.8) e il successivo esempio.

Appendice A

Algebra lineare

A.1 polinomi

Teorema A.1.1 (di Cayley-Hamilton). *Si consideri la matrice $A \in \mathbb{R}^{n \times n}$. Si ha, allora, $\Delta_A(A) = 0$.*

Il risultato è valido anche per $A \in \mathbb{K}^{n \times n}$, dove $\mathbb{K}$ è un campo. In particolare vale per $\mathbb{K} = \mathbb{C}$.

La dimostrazione utilizza lo strumento delle matrici polinomiali, ovvero matrici i cui elementi sono polinomi in $K(s)$.

Dimostrazione A.1.1. Si consideri il polinomio caratteristico di A:

$$\Delta_A(s) = \det(sI - A) = s^n + \alpha_{n-1}s^{n-1} + \cdots + \alpha_0 \quad (\text{A.1})$$

$$L_0(s) = 0$$

$$L_i(s) = A^{i-1} + \cdots + s^{i-1}I \quad i = 1, 2, \dots, n.$$

Sostituendo A^i con $s^i I - L_i(s)(sI - A)$ si ottiene

$$\begin{aligned} & \sum_{i=0}^n \alpha_i s^i I - \sum_{i=0}^n L_i(s)(sI - A)\alpha_i = \\ & = \Delta_A(s)I - \sum_{i=0}^n L_i(s)(sI - A)\alpha_i = \\ & = \text{adj}(sI - A)(sI - A) - \sum_{i=0}^n L_i(s)(sI - A)\alpha_i = \\ & = (\text{adj}(sI - A) - \sum_{i=0}^n L_i(s)\alpha_i)(sI - A) = \\ & = M(s)(sI - A) \end{aligned}$$

Definizione A.1.1 (Matrice Singolare). Sia data una matrice $A \in \mathbb{R}^{n \times n}$. La matrice A si dice singolare se il suo determinante è nullo, cioè

$$\Delta(A) = 0 \quad (\text{A.2})$$

Definizione A.1.2 (Matrice dei valori singolari di Sylvester). Sia A una generica matrice, $A \in \mathbb{R}^{m \times n}$. Sia $A^\dagger = A \cdot A^\top$ la matrice quadrata ottenuta moltiplicando la matrice per la sua trasposta. Si dice *Matrice dei valori singolari di Sylvester* la matrice $S \in \mathbb{R}^{m \times n}$ nella forma

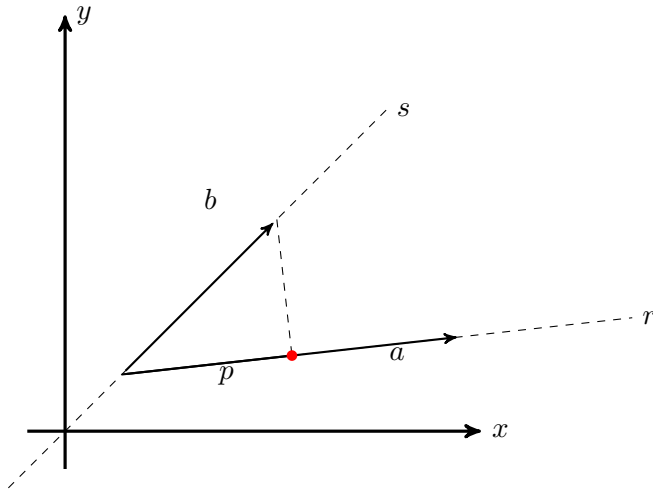
$$S \triangleq \begin{bmatrix} s_{11} & 0 & 0 & 0 & \dots \\ 0 & s_{22} & 0 & 0 & \dots \\ 0 & 0 & s_{ij} & 0 & \dots \\ 0 & 0 & 0 & \ddots & \dots \end{bmatrix} \quad \text{oppure} \quad S = \begin{bmatrix} s_{11} & 0 & 0 & 0 \\ 0 & s_{22} & 0 & 0 \\ 0 & 0 & s_{ij} & 0 \\ 0 & 0 & 0 & \ddots \\ \vdots & \vdots & \vdots & \ddots \end{bmatrix} \quad (\text{A.3})$$

i cui coefficienti non nulli, denotati con s_{ij} , sono presenti solo sulla pseudo-diagonale. Tali coefficienti sono detti *valori singolari* e si ottengono estraendo le radici quadrate degli autovalori di $A^\dagger$.

Si noti che la matrice dei valori singolari non è, in generale, quadrata: le sue dimensioni dipendono dalle esigenze, in particolare dalle dimensioni delle matrici coinvolte nell'algoritmo S.V.D., in cui essa è maggiormente utilizzata.

A.2 Proiezioni ortogonali

Si consideri una retta direttrice r nel piano a due dimensioni, $r \in \mathbb{R}^2$ nel contesto di un sistema di riferimento ortonormale. Sia a un vettore appartenente a tale retta.



Si consideri una seconda retta $s \in \mathbb{R}^2$, complanare. Sia b un vettore appartenente a tale retta. Si vuole determinare un'espressione per il vettore p , cioè la proiezione di b su a secondo la direttrice ortogonale a r . Si suppone che p sia legato ad a secondo una relazione di proporzionalità:

$$\mathbf{p} = \mathbf{a}x$$

si vuole quindi che a sia ortogonale al vettore differenza $c = b - p$, cioè che

$$\mathbf{a}^\top \cdot (\mathbf{b} - \mathbf{p}) = 0$$

$$\mathbf{a}^\top \cdot (\mathbf{b} - \mathbf{a}x) = 0$$

Sviluppando si ottiene:

$$\mathbf{a}^\top \mathbf{a}x = \mathbf{a}^\top \mathbf{b}$$

E, finalmente:

$$x = \frac{\mathbf{a}^\top \mathbf{b}}{\mathbf{a}^\top \mathbf{a}}$$

quindi il vettore $\mathbf{p}$ è univocamente determinato:

$$\mathbf{p} = \mathbf{a}x = \mathbf{a} \frac{\mathbf{a}^\top \mathbf{b}}{\mathbf{a}^\top \mathbf{a}} = \mathbf{P}\mathbf{b} \quad (\text{A.5})$$

La componente

$$\mathbf{P} = \mathbf{a} \frac{\mathbf{a}^\top}{\mathbf{a}^\top \mathbf{a}} \quad (\text{A.6})$$

è detta *matrice di proiezione*, ed è il legame cercato tra un vettore $\mathbf{b}$ e la sua proiezione $\mathbf{p}$. Se A è definito su uno spazio vettoriale su più dimensioni (cioè può essere rappresentato mediante una matrice) esiste una formula analoga che si può ricavare mediante lo stesso procedimento

$$A^\top AX = A^\top \mathbf{b}$$

La matrice X è incognita; si può ricavare risolvendo il sistema. La soluzione è

$$X = (A^\top A)^{-1} A^\top \mathbf{b}$$

La proiezione ortogonale $\mathbf{p}$ vale

$$\mathbf{p} = AX = A(A^\top A)^{-1} A^\top \mathbf{b} = \mathbf{P}\mathbf{b} \quad (\text{A.7})$$

e finalmente si può formulare la seguente

Definizione A.2.1. La matrice di proiezione $\mathbf{P}$ è la matrice $\in \mathbb{R}^{n \times n}$ definita da

$$\mathbf{P} = A(A^\top A)^{-1} A^\top \quad (\text{A.8})$$

Proprietà A.2.0.1. La matrice di proiezione $\mathbf{P}$ ha lo stesso spazio delle colonne di A : $\text{Im}(\mathbf{P}) = \text{Im}(A)$, cioè $\text{Im}(\mathbf{P}) \subseteq \text{Im}(A)$ e $\text{Im}(A) \subseteq \text{Im}(\mathbf{P})$.

Proprietà A.2.0.2. La matrice di proiezione è idempotente, cioè $\mathbf{P} = \mathbf{P}^2$

Dimostrazione A.2.1. Esplicitando il prodotto $\mathbf{P}^2 = \mathbf{P}\mathbf{P}$ si ottiene:

$$\mathbf{P}\mathbf{P} = A(A^\top A)^{-1} \underbrace{A^\top A(A^\top A)^{-1}}_{=I} A^\top = A(A^\top A)^{-1} A^\top$$

Se la matrice $A^\top A$ è invertibile è possibile parlare di proiezione ortogonale. In questo caso la matrice gode della seguente

Proprietà A.2.0.3. La matrice di proiezione $\mathbf{P}$ è simmetrica, cioè $\mathbf{P} = \mathbf{P}^\top$.

Dimostrazione A.2.2. Se la matrice $\mathcal{A} = A^\top A$ è invertibile, essa gode della identità

$$[\mathcal{A}^\top]^{-1} = [\mathcal{A}^{-1}]^\top$$

e inoltre essa sarà simmetrica, ovvero $\mathcal{A} = \mathcal{A}^\top$. Perciò sarà possibile scrivere la catena di uguaglianze:

$$\mathbf{P}^\top = (A^\top)^\top \left((A^\top A)^{-1} \right)^\top A^\top = A \left((A^\top A)^\top \right)^{-1} A^\top = A(A^\top A)^{-1} A^\top = \mathbf{P}$$

Osservazione 3. Nel caso $A^\top A$ non sia invertibile la matrice di proiezione, $\mathbf{P}^*$, prende il nome di *operatore di proiezione* e si può calcolare sfruttando la pseudo-inversa di Moore-Penrose. Essa assumerà una espressione formalmente identica a (A.8). In questo caso però la proiezione che si otterrà non sarà ortogonale:

$$P^* = A(A^\top A)^\# A^\top$$

A.3 Matrici e sottospazi

Per un qualunque isomorfismo descritto da una matrice $A \in \mathbb{R}^{m \times n}$ si possono individuare quattro *sottospazi fondamentali*. Essi sono

$$\mathcal{N}(A), \ker(A) \quad \text{Lo spazio nullo, o nullità. E' indicato anche con nucleo, o kernel} \quad (\text{A.9a})$$

$$\mathcal{R}(A), \text{Im}(A) \quad \text{Lo spazio delle colonne. E' possibile indicarlo con Immagine} \quad (\text{A.9b})$$

$$\mathcal{N}(A^\top) \quad \text{Lo spazio nullo sinistro, ovvero il nucleo della trasposta.} \quad (\text{A.9c})$$

$$\mathcal{R}(A^\top) \quad \text{Lo spazio delle righe, ovvero le colonne della trasposta.} \quad (\text{A.9d})$$

Tra i sottospazi di A sussistono alcune relazioni fondamentali.

$$\mathcal{R}(A) = [\mathcal{N}(A^\top)]^\perp \quad \mathcal{R}(A^\top) = [\mathcal{N}(A)]^\perp \quad (\text{A.10a})$$

$$\mathcal{N}(A) = [\mathcal{R}(A^\top)]^\perp \quad \mathcal{N}(A^\top) = [\mathcal{R}(A)]^\perp \quad (\text{A.10b})$$

Tali relazioni costituiscono una forma dei cosiddetti teoremi fondamentali dell'algebra.

Dimostrazione A.3.1. Se un vettore v appartiene al nucleo, cioè $v \in \ker A$ si può scrivere che:

$$v \in \ker A \iff Av = 0 \iff \underbrace{v^\top A^\top}_{\text{per le proprietà della trasposta}} = 0 \iff v^\top A^\top w = 0, \forall w \in \mathbb{R}^m$$

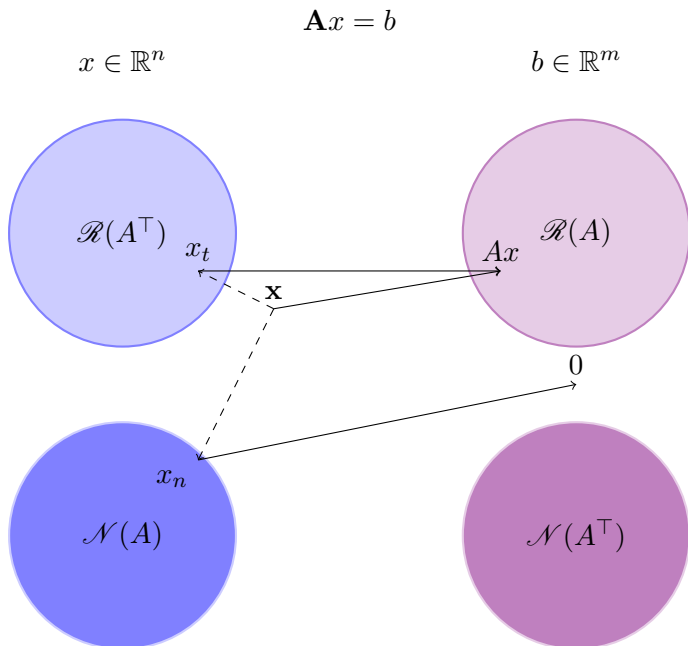
e quindi $v \in [\text{Im}(A^\top)]^\perp$. Se nucleo e immagine sono l'uno il complemento ortogonale dell'altro vale anche la prima.

Se A è una applicazione lineare, cioè una scrittura della forma

$$A\mathbf{x} = \mathbf{b} \quad (\text{A.11})$$

Dove $A \in \mathbb{R}^{m \times n}$, $x \in \mathbf{X} \subset \mathbb{R}^n$ e $b \in \mathbf{B} \subset \mathbb{R}^m$.

I legami tra i quattro sottospazi si possono rappresentare secondo i seguenti diagrammi i quali descrivono -in sintesi- l'effetto di un generico isomorfismo:



Un generico vettore $\mathbf{x}$ viene scomposto in due componenti:

1. x_n , appartenente allo spazio nullo e
2. x_t , appartenente allo spazio delle righe.

La prima verrà portata in 0, la seconda nel termine $b \in \mathbb{R}^m$.

A.4 Pseudoinversa di Moore-Penrose

A.4.1 definizioni

Nel contesto di questa analisi si introduce una possibile trasformazione, applicabile a matrici non quadrate, da intendersi come una generalizzazione della *inversa* di una matrice.

Definizione A.4.1 (inversa di Moore-Penrose). Nel contesto di un generico corpo commutativo (o campo), si definisca una matrice $A \in \mathbb{K}^{m \times n}$. La pseudoinversa di Moore-Penrose di A è definita come la matrice $A^\# \in \mathbb{K}^{n \times m}$ che soddisfa tutte le seguenti condizioni:¹

$$AA^\#A = A \quad AA^\# \text{ deve mandare tutti i vettori colonna di } A \text{ in se stessi.} \quad (\text{A.12a})$$

$$A^\#AA^\# = A^\# \quad A \text{ è un inversa debole per il semigrupp moltiplicativo.} \quad (\text{A.12b})$$

$$(AA^\#)^* = AA^\# \quad AA^\# \text{ è Hermitiana.} \quad (\text{A.12c})$$

$$(A^\#A)^* = A^\#A \quad A^\#A \text{ è Hermitiana.} \quad (\text{A.12d})$$

Nel caso particolare in cui il campo $\mathbb{K} = \mathbb{R}$ la (A.12c) e la (A.12d) devono essere intese come segue:

$$(AA^\#)^\top = AA^\# \quad AA^\# \text{ è uguale alla sua trasposta.} \quad (\text{A.12e})$$

$$(A^\#A)^\top = A^\#A \quad A^\#A \text{ è uguale alla sua trasposta.} \quad (\text{A.12f})$$

La pseudoinversa di Moore-Penrose rappresenta una classe di trasformazioni che è solo un sotto-insieme rispetto alla classe delle cosiddette inverse generalizzate.

Tuttavia tale trasformazione ha alcune interessanti proprietà.

Proprietà A.4.1.1. Se i coefficienti di A sono reali lo saranno anche quelli di $A^\#$: $a_{ij} \in \mathbb{R} \Leftrightarrow a_{ij}^\# \in \mathbb{R}$.

Proprietà A.4.1.2. Se A è invertibile, la pseudoinversa e l'inversa coincidono: $A^\# = A^{-1}$.

Proprietà A.4.1.3. La pseudoinversa di una pseudoinversa è la matrice originale: $(A^\#)^\# = A$.

Proprietà A.4.1.4. La pseudoinversa gode della proprietà commutativa nei riguardi delle seguenti operazioni:

$$\text{trasposizione:} \quad (A^\top)^\# = (A^\#)^\top \quad (\text{A.13a})$$

$$\text{coniugazione:} \quad (\overline{A})^\# = \overline{A^\#} \quad (\text{A.13b})$$

$$\text{trasposizione Hermitiana:} \quad (A^*)^\# = (A^\#)^* \quad (\text{A.13c})$$

¹Si userà la notazione $A^\top$ per indicare la trasposta e A^* per indicare la trasposta Hermitiana.

Proprietà A.4.1.5. La pseudoinversa si può scrivere come segue:

$$A^\# = (A^\top \cdot A)^\# A^\top = A^\top (A \cdot A^\top)^\# \quad (\text{A.14})$$

Se $A \in \mathbb{R}^{m \times n}$ e se $m < n$ e il rango di A $\text{rank}(A) = m$, allora $A \cdot A^\top$ è non singolare e si può scrivere la pseudoinversa come

$$A^\# = A^\top (AA^\top)^{-1} \quad (\text{A.15})$$

La quale viene chiamata *pseudoinversa destra*. Viceversa, se $n < m$ e il rango di A $\text{rank}(A) = n$, allora $A^\top \cdot A$ è non singolare e la pseudoinversa diviene

$$A^\# = (A^\top A)^{-1} A^\top \quad (\text{A.16})$$

Che prende il nome di *pseudoinversa sinistra*.

Definizione A.4.2 (Proiettori). Le matrici

$$P = AA^\# \quad (\text{A.17a})$$

$$Q = A^\# A \quad (\text{A.17b})$$

identificano speciali operatori.

- P , nella (A.17a) è detto proiettore ortogonale su $\mathcal{R}(A)$, cioè sullo spazio generato dalle colonne di A .

Proprietà A.4.1.6. $\text{Im}(A) = \text{Im}(AA^\#)$.

La definizione di pseudoinversa garantisce che $\text{Im}(AA^\#) \subseteq \text{Im}(A)$.

Inoltre, si ha $\text{Im}(A) = \text{Im}(AA^\#A) \subseteq \text{Im}(AA^\#)$.

- Q , nella (A.17b) è detto proiettore ortogonale su $\mathcal{R}(A^\top)$, ovvero sullo spazio generato dalle colonne della sua trasposta, quindi le righe.

Proprietà A.4.1.7. $\text{Im}(A^\top) = \text{Im}(A^\#A) = \text{Im}(A^\#)$.

Ciascuno dei proiettori obbedisce al seguente

Teorema A.4.1. Se $\mathcal{P}$ è un proiettore ortogonale nello spazio $\mathcal{V}$ allora $I - \mathcal{P}$ è un proiettore ortogonale nello spazio $\mathcal{V}^\perp$.

Dimostrazione A.4.1. $\forall x, y \in \mathbb{R}^n$ è possibile scegliere $v \in \mathcal{V}$ e $w \perp v$. Sarà sufficiente fissare $v := \mathcal{P}x$ e $w := (I - \mathcal{P})y$. Da ciò discende che

$$v^\top w = x^\top \mathcal{P}(I - \mathcal{P})y = 0 \Rightarrow \text{Im}(I - \mathcal{P}) \subseteq \mathcal{V}^\perp$$

Per converso, se si suppone che $x \in \mathcal{V}^\perp$ si ha:

$$0 = x^\top \mathcal{P}x = \underbrace{x^\top \mathcal{P} \mathcal{P} x}_{\text{idempotenza delle matrici di proiezione}} = x^\top \mathcal{P}^\top \mathcal{P} x = ||\mathcal{P}||^2 x = 0$$

Quindi

$$\Rightarrow I - \mathcal{P}x = x \Rightarrow x \in \text{Im}(I - \mathcal{P})$$

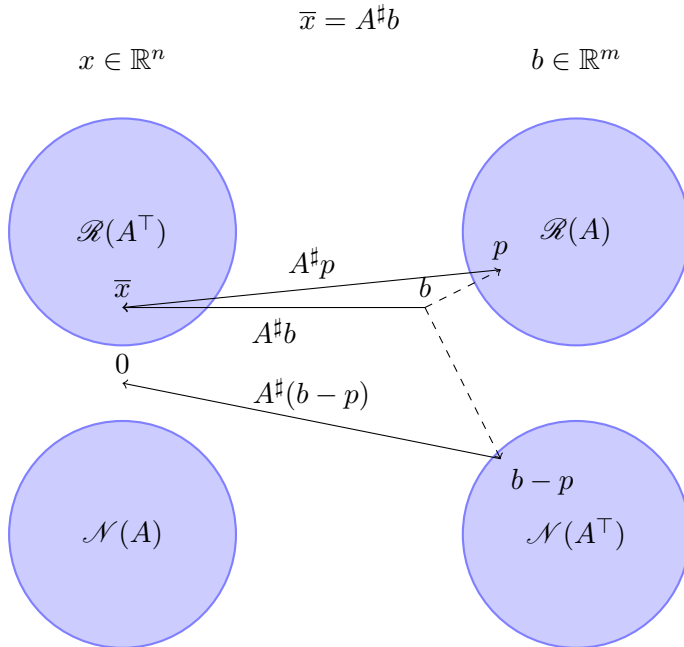
I proiettori godono delle seguenti proprietà:

Proprietà A.4.1.8. $(I - P)$ è il proiettore ortogonale in $\mathcal{N}(A^\top)$, cioè nel nucleo $\ker(A^\top)$.

Proprietà A.4.1.9. $(I - Q)$ è il proiettore ortogonale nel nucleo di A , $\mathcal{N}(A)$.

A.4.2 Relazioni tra i sottospazi

L'azione della pseudoinversa $A^\sharp$ si può agevolmente rappresentare mediante un diagramma converso, come il seguente.



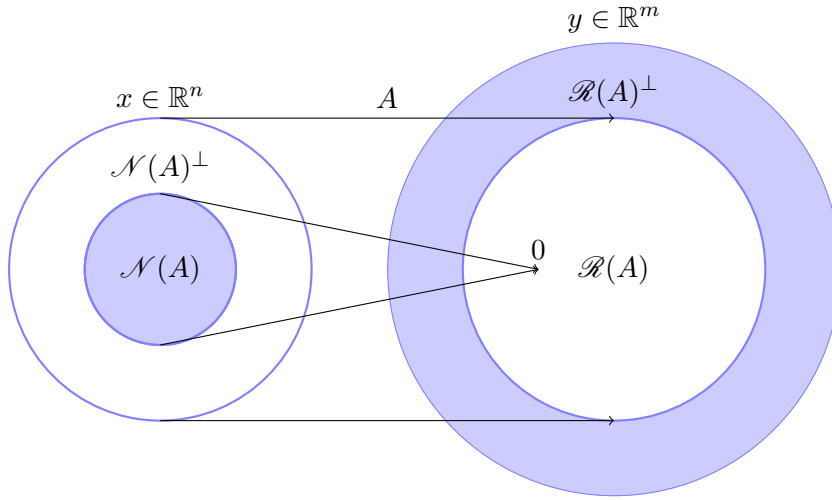
Nel quale si intuisce che l'effetto della pseudoinversa può esplicarsi secondo due passi distinti:

1. Il vettore b verrà scomposto in componenti ortogonali, ovvero: $b - p$, appartenente allo spazio nullo sinistro e p , appartenente allo spazio delle colonne.
2. La componente dello spazio nullo verrà mandata in 0, l'altra in $\bar{x}$ che risolve il sistema.

Se si definiscono ora i sottospazi seguenti:

$$\begin{array}{ll} \mathcal{R}(A)^\perp & \text{per indicare il complemento ortogonale di } \mathcal{R}(A) \\ \mathcal{N}(A)^\perp & \text{per indicare il complemento ortogonale di } \mathcal{N}(A) \end{array}$$

che si rappresentano facilmente mediante i seguenti diagrammi di Eulero-Venn nella pagina seguente:



Le proprietà della pseudoinversa $A^\sharp$ giustificano la scrittura di ulteriori relazioni tra i sottospazi appena descritti.

$$\mathcal{R}(A) = \mathcal{N}(A^\sharp)^\perp = \mathcal{R}(AA^\sharp) = \mathcal{N}(I - AA^\sharp) \quad (\text{A.18a})$$

$$\mathcal{R}(A)^\perp = \mathcal{N}(A^\sharp) = \mathcal{N}(AA^\sharp) = \mathcal{R}(I - AA^\sharp) \quad (\text{A.18b})$$

$$\mathcal{N}(A) = \mathcal{R}(A^\sharp)^\perp = \mathcal{N}(A^\sharp A) = \mathcal{R}(I - A^\sharp A) \quad (\text{A.18c})$$

$$\mathcal{N}(A)^\perp = \mathcal{R}(A^\sharp) = \mathcal{R}(A^\sharp A) = \mathcal{N}(I - A^\sharp A) \quad (\text{A.18d})$$

Se A non è invertibile, e il sistema lineare associato è sottodeterminato ($m < n$) le proprietà precedenti consentono di stabilire che tutte le soluzioni del sistema (A.11) avranno la forma:

$$\mathbf{x} = A^\sharp b + [I - A^\sharp A]\mathbf{w} \quad (\text{A.19})$$

con $\mathbf{w}$ vettore arbitrario. Le soluzioni esistono se e solo se $AA^\sharp b = b$.

se quest'ultima è valida, la soluzione è unica se e solo se la matrice A ha rango pieno. In tale caso la matrice $[I - A^\sharp A]$ è una matrice nulla.

A.5 Esempi

A.5.1 Calcolo della Pseudoinversa

Esempio A.5.1.1. Se la matrice è non quadrata ma ha rango *per righe* o *per colonne* pieno, si potranno sfruttare le proprietà (A.15) e (A.16).

Si consideri, come esempio, la matrice

$$B = \begin{bmatrix} 1 & -1 & 0 \\ 0 & 1 & 1 \end{bmatrix}$$

A partire dalle proprietà citate si ottiene:

$$\begin{aligned}
B^\sharp &= B^T(BB^T)^{-1} \\
&= \begin{bmatrix} 1 & 0 \\ -1 & 1 \\ 0 & 1 \end{bmatrix} \left\{ \begin{bmatrix} 1 & -1 & 0 \\ 0 & 1 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ -1 & 1 \\ 0 & 1 \end{bmatrix} \right\}^{-1} \\
&= \begin{bmatrix} 1 & 0 \\ -1 & 1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix}^{-1} \\
&= \frac{1}{3} \begin{bmatrix} 2 & 1 \\ -1 & 1 \\ 1 & 2 \end{bmatrix}
\end{aligned}$$

Esempio A.5.1.2 (decomposizione S.V.D. singular value decomposition). Se la matrice non ha rango pieno non si possono usare le proprietà sopra esposte. Il calcolo della pseudo-inversa di una matrice A può essere eseguito sfruttando una decomposizione, ovvero ottenendo A con un prodotto di tre matrici con particolari proprietà.

La decomposizione in valori singolari secondo Sylvester (1889) permette di scrivere una qualunque matrice A come un prodotto del tipo

$$A = USV^\top \quad (\text{A.20})$$

Dove la matrice S è la *matrice dei valori singolari* definita precedentemente e V e U sono matrici ortonormali. A titolo d'esempio, si consideri la matrice

$$A = \begin{bmatrix} 4 & 0 & 2 & 6 \\ 2 & 0 & 1 & 3 \end{bmatrix} \quad \mathcal{L}_A : \mathbb{R}^4 \rightarrow \mathbb{R}^2$$

Nella quale, evidentemente, le due righe sono proporzionali e quindi non indipendenti. Il rango è $\text{rk}(A) = 1$. La matrice V si può scrivere ottenendo lo spazio vettoriale del dominio $\mathbb{R}^4$ come *somma diretta* di due sottospazi, generati da due basi ortogonali.

$$\begin{aligned}
\mathbb{R}^4 = & \underbrace{\left\langle \begin{bmatrix} 2 \\ 0 \\ 1 \\ 3 \end{bmatrix} \right\rangle}_{=\mathcal{C}(A^\top)=\text{Im}(A^\top)} \oplus \underbrace{\left\langle \begin{bmatrix} 0 \\ 1 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} -1 \\ 0 \\ 2 \\ 0 \end{bmatrix}, \begin{bmatrix} -3 \\ 0 \\ 0 \\ 2 \end{bmatrix} \right\rangle}_{=\mathcal{N}(A)=\text{ker}(A)}
\end{aligned}$$

Ora i vettori della base di $\mathcal{N}(A)$, cioè il nucleo di A non sono ortogonali tra di loro. Se si procede con una opportuna ortonormalizzazione è possibile ottenere una matrice V ortonormale che contenga le basi di $\mathcal{C}(A^\top)$ e $\mathcal{N}(A)$. La matrice V così configurata potrà finalmente essere scritta nella forma

$$V = \begin{bmatrix} \frac{2}{\sqrt{14}} & 0 & 0 & \frac{5}{\sqrt{35}} \\ 0 & 1 & 0 & 0 \\ \frac{1}{\sqrt{14}} & 0 & \frac{3}{\sqrt{10}} & \frac{-1}{\sqrt{35}} \\ \frac{3}{\sqrt{14}} & 0 & \frac{-1}{\sqrt{10}} & \frac{-3}{\sqrt{35}} \end{bmatrix} = [v_1 \quad v_2 \quad v_3 \quad v_4]$$

Per ciò che riguarda il codominio si potrà scrivere, allo stesso modo:

$$\begin{aligned}
\mathbb{R}^2 = & \underbrace{\left\langle \begin{bmatrix} 2 \\ 1 \end{bmatrix} \right\rangle}_{=\mathcal{C}(A)=\text{Im}(A)} \oplus \underbrace{\left\langle \begin{bmatrix} 1 \\ -2 \end{bmatrix} \right\rangle}_{=\mathcal{N}(A^\top)=\text{ker } A^\top}
\end{aligned}$$

e si potrà ottenere una matrice ortonormale, che avrà la forma:

$$U = \begin{bmatrix} \frac{2}{\sqrt{5}} & \frac{1}{\sqrt{5}} \\ \frac{1}{\sqrt{5}} & \frac{-2}{\sqrt{5}} \end{bmatrix} = [u_1 \quad u_2]$$

La matrice $A^\perp = A \cdot A^\top$ è:

$$A^\perp = \begin{bmatrix} 56 & 28 \\ 28 & 14 \end{bmatrix}$$

Il cui polinomio caratteristico è

$$p(A^\perp) = (56 - \lambda)(14 - \lambda) - (784) = (\lambda)^2 - 70\lambda = \lambda(\lambda - 70)$$

e i cui autovalori sono $a_1^\perp = 0$ e $a_2^\perp = 70$.

Quindi la matrice dei valori singolari si potrà costruire come:

$$S = \begin{bmatrix} \sqrt{70} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

Mentre la matrice A si potrà decomporre nella forma (A.20):

$$A = \begin{bmatrix} \frac{2}{\sqrt{5}} & \frac{1}{\sqrt{5}} \\ \frac{1}{\sqrt{5}} & \frac{-2}{\sqrt{5}} \end{bmatrix} \cdot \begin{bmatrix} \sqrt{70} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} \cdot \begin{bmatrix} \frac{2}{\sqrt{14}} & 0 & \frac{1}{\sqrt{14}} & \frac{3}{\sqrt{14}} \\ 0 & 1 & 0 & 0 \\ 0 & 0 & \frac{3}{\sqrt{10}} & \frac{-1}{\sqrt{10}} \\ \frac{5}{\sqrt{35}} & 0 & \frac{\sqrt{10}}{\sqrt{35}} & \frac{-3}{\sqrt{35}} \end{bmatrix}$$

e finalmente la pseudo-inversa di Moore-Penrose potrà essere calcolata come segue:

$$A^\# = V \cdot S^\# \cdot U^\top$$

con $S^\#$ la pseudoinversa di S . In questo caso S è una matrice nella forma indicata in (A.3), ed $S^\#$ si può ottenere in modo diretto trasponendo S e sostituendone gli elementi non nulli con i loro reciproci. Quindi si potrà scrivere:

$$A^\# = \begin{bmatrix} \frac{2}{\sqrt{14}} & 0 & 0 & \frac{5}{\sqrt{35}} \\ 0 & 1 & 0 & 0 \\ \frac{1}{\sqrt{14}} & 0 & \frac{3}{\sqrt{10}} & \frac{-1}{\sqrt{35}} \\ \frac{3}{\sqrt{14}} & 0 & \frac{-1}{\sqrt{10}} & \frac{-3}{\sqrt{35}} \end{bmatrix} \begin{bmatrix} \frac{1}{\sqrt{70}} & 0 \\ 0 & 0 \\ 0 & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} \frac{2}{\sqrt{5}} & \frac{1}{\sqrt{5}} \\ \frac{1}{\sqrt{5}} & \frac{-2}{\sqrt{5}} \end{bmatrix} = \frac{1}{70} \begin{bmatrix} 4 & 2 \\ 0 & 0 \\ 2 & 1 \\ 6 & 3 \end{bmatrix}$$

Il metodo S.V.D è quindi efficace anche se le matrici non hanno rango pieno. Rappresenta perciò un metodo generale.

Esempio A.5.1.3 (decomposizione LU). E' possibile individuare un terzo metodo, che fa uso della decomposizione LU . La decomposizione LU è possibile per qualsiasi matrice C , che può essere scritta come il prodotto

$$C = LU$$

con U una matrice triangolare superiore, L una triangolare inferiore. L'algoritmo di decomposizione consiste nel ridurre in forma canonica la matrice A *tenendo traccia* delle operazioni eseguite (ovvero delle matrici di passaggio).

Si consideri la matrice

$$C = \begin{bmatrix} 1 & 2 \\ 2 & 5 \\ 3 & 7 \end{bmatrix}$$

Si osserva subito che la terza riga è il risultato della somma delle prime due. La terza riga quindi è destinata ad annullarsi:

$$C = \underbrace{\begin{bmatrix} 1 & 2 \\ 2 & 5 \\ 3 & 7 \end{bmatrix}}_{E_1} \rightsquigarrow \begin{bmatrix} 1 & 2 \\ 2 & 5 \\ 0 & 0 \end{bmatrix}$$

$$E_1 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ -1 & -1 & 1 \end{pmatrix}$$

Con un successivo passaggio si può sottrarre dalla seconda riga il doppio della prima, per ottenere un pivot:

$$C = \underbrace{\begin{bmatrix} 1 & 2 \\ 2 & 5 \\ 0 & 0 \end{bmatrix}}_{E_2} \rightsquigarrow \begin{bmatrix} 1 & 2 \\ 0 & 1 \\ 0 & 0 \end{bmatrix}$$

$$E_2 = \begin{pmatrix} 1 & 0 & 0 \\ -2 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Quindi le matrici U e L sono univocamente determinate:

$$U = \begin{bmatrix} 1 & 2 \\ 0 & 1 \\ 0 & 0 \end{bmatrix} \quad L = (E_1)^{-1}(E_2)^{-1} = \begin{bmatrix} 1 & 0 & 0 \\ 2 & 1 & 0 \\ 3 & 1 & 1 \end{bmatrix}$$

Si noti che il prodotto LU non cambia se si cancella l'ultima riga di U e l'ultima colonna di L , il che può semplificare il successivo calcolo. Rinomineremo le nuove matrici come segue:

$$\bar{U} = \begin{bmatrix} 1 & 2 \\ 0 & 1 \end{bmatrix} \quad \bar{L} = \begin{bmatrix} 1 & 0 \\ 2 & 1 \\ 3 & 1 \end{bmatrix}$$

Grazie alla fattorizzazione LU è possibile arrivare ad una definizione costruttiva della matrice pseudoinversa $C^\sharp$.

$$C^\sharp = \bar{U}^T (\bar{U}\bar{U}^T)^{-1} (\bar{L}^T \bar{L})^{-1} \bar{L}^T \quad (\text{A.21})$$

Finalmente si potrà arrivare al risultato:

$$C^\sharp = \begin{bmatrix} 1 & 0 \\ 2 & 1 \end{bmatrix} \begin{bmatrix} 1 & -2 \\ -2 & 5 \end{bmatrix} \frac{1}{3} \begin{bmatrix} 2 & -5 \\ -5 & 14 \end{bmatrix} \begin{bmatrix} 1 & 2 & 3 \\ 0 & 1 & 1 \end{bmatrix} = \begin{bmatrix} -4 & -3 & 1 \\ -\frac{5}{3} & \frac{4}{3} & -\frac{1}{3} \end{bmatrix}$$

A.5.2 calcolo della matrice di proiezione

Esempio A.5.2.1. Trovare la matrice di proiezione $\mathbf{P}$ sul piano $\in \mathbb{R}^3$ definito da

$$\alpha : x + y - z = 0$$

Ora il vettore direttore, cioè il vettore che risulta ortogonale al piano è definito da

$$N = \begin{pmatrix} 1 \\ 1 \\ -1 \end{pmatrix}$$

Il sottospazio sul quale si sta proiettando è definito dalla matrice A , le cui colonne devono contenere le basi che generano il sottospazio definito dal piano α . Essendo in tre dimensioni saranno sufficienti due vettori per specificare univocamente il piano. Una scelta non unica per i vettori che soddisfatti l'equazione α è

$$a_1 = \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix} \qquad a_2 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$$

allora la matrice A può essere descritta da

$$A = \begin{bmatrix} 1 & 1 \\ -1 & 0 \\ 0 & 1 \end{bmatrix}$$

La matrice di proiezione si potrà facilmente ricavare dalla (A.8).

$$\mathbf{P} = \begin{bmatrix} 1 & 1 \\ -1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} 2 & 1 \\ 1 & 2 \end{bmatrix}^{-1} \begin{bmatrix} 1 & -1 & 0 \\ 1 & 0 & 1 \end{bmatrix} = \dots = \frac{1}{3} \begin{bmatrix} 2 & -1 & 1 \\ -1 & 2 & 1 \\ 1 & 1 & 2 \end{bmatrix}$$

Bibliografia

- [1] Ferrante, Augusto and Wimmer, K. Harald and Minghua, Lin (2011), *Block companion matrices and reachability matrices*, Electronic Journal of Linear Algebra.
- [2] Ferrante, Augusto and Wimmer, K. Harald (2010), *Reachability matrices and cyclic matrices*, Electron. J. Linear Algebra 20, 95-102.
- [3] Strang, Gilbert (1981), *Linear Algebra and applications*, Wellesley-Cambridge press.
- [4] Nakamura, Yoshihiko and Takano, Wataru (1991), *Advanced Robotics*, Addison-Wesley Longman, Incorporated.
- [5] Cassandras, Christos G. and LaFortune, Stéphane (2008), *Introduction to Discrete Event systems*, Springer.