

Periodico di Matematiche



**Organo della
MATHESIS**

Società italiana di scienze
matematiche e fisiche
fondata nel 1895

**Volume 100 (2024)
Fascicolo 1/2**

Direttore

Salvatore Cuomo

salvatore.cuomo@unina.it

Comitato di Redazione

Maria Coccozza

mariacoccozza@lscortese.com

Atalia Del Bene

ataliadelbene1960@gmail.com

Umberto Dello Iacono

umberto.delloiacono@unicampania.it

Domenica Di Sorbo

presidente-disorbo@mathesisnazionale.it

Lorenzo Mazza

lorenzomazza@gmail.com

Marcello Pedone

marcellopedone@tin.it

Alessio Russo

alessio.russo@unicampania.it

Annalisa Santini

annalisasantini66@gmail.com

Francesco Sicolo

fsicolo@gmail.com

Autorizzazioni e supporti Autorizzazioni Tribunale di Bologna n. 266 del 29-3-1950

Il Periodico di Matematiche è distribuito gratuitamente ai soci Mathesis. Le istruzioni per associarsi su www.mathesisnazionale.it

Abbonamenti per Scuole ed Enti Vari: Per l'Italia €60,00 - Per l'Estero €70,00

Per richiesta di numeri singoli o arretrati: segreteria@mathesisnazionale.it

c/c postale, Codice IBAN:

IT05I0760104000000048597470

intestato a:

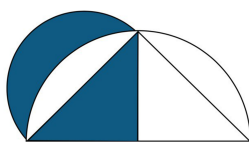
Mathesis Nazionale c/o Dipartimento di Matematica e Fisica
Università degli Studi della Campania "Luigi Vanvitelli"
Viale Abramo Lincoln 5 - 81100 Caserta (CE)
www.mathesisnazionale.it • info@mathesisnazionale.it

Volume **100** Fascicolo 1-2 (2024)

Periodico di Matematiche

Organo della MATHESIS

*Società italiana di scienze
fisiche e matematiche
fondata nel 1895*



Mathesis

Indice

Salvatore Cuomo e Domenica Di Sorbo

Editoriale dedicato allo scomparso prof. Francesco de Giovanni i

Lorenzo Millanti

Indice di correlazione generalizzato, un'alternativa a Kendall 1

Franco Fericola

Circocentro e Ortocentro di un triangolo 11

Nicola Fusco

Crittografia: una Prospettiva Didattica – parte 1 21

Antonino Onorato

L'equazione differenziale di Karl Pearson 35

Giovanni Di Mare e Domenico Veca

Le "infinite" identità del fattoriale, dai reali agli ipercomplessi 63

Letizia Pellegrini e Alberto Peretti

Una procedura per risolvere un sistema lineare di complementarità 83

Corrado Binetti

Proprietà elastiche e problemi di misura ottimale: un approccio didattico 95

Luca Granieri

Matematica quantistica 103

Nicola Fusco

Crittografia: una Prospettiva Didattica – parte 2 121

Editoriale

dedicato allo scomparso prof. Francesco de Giovanni

Con profondo dolore e sincera commozione, questo editoriale si apre nel ricordo del compianto Direttore della rivista, il Professor Francesco de Giovanni, una figura di rilievo nel panorama matematico italiano e internazionale. Già Presidente della Mathesis Nazionale e Professore Ordinario di Algebra presso l'Università degli Studi di Napoli Federico II, il Professor de Giovanni ha lasciato un'impronta indelebile nella comunità scientifica e accademica.

Nato a Napoli nel 1955, Francesco de Giovanni si è laureato in Matematica nel 1978 presso l'Università degli Studi di Napoli Federico II, sotto la guida del Professor Mario Curzio. La sua formazione si è poi arricchita grazie al periodo di studi presso l'Università dell'Illinois a Urbana-Champaign (USA), sotto la supervisione del Professor D.J.S. Robinson. Nel 1987, è stato nominato Professore Ordinario di Algebra presso l'ateneo partenopeo, ruolo che ha ricoperto con passione e dedizione fino agli ultimi anni della sua carriera.

Il Professor de Giovanni ha contribuito in maniera significativa allo sviluppo della teoria dei gruppi infiniti, un ambito nel quale ha pubblicato oltre 250 articoli scientifici. I suoi studi si sono concentrati su temi fondamentali quali le proprietà dei gruppi fattorizzati, le questioni reticolari, le condizioni finitarie, il comportamento degli automorfismi di gruppi, i gruppi di rango infinito, i gruppi di elevata cardinalità e i gruppi lineari infiniti. La profondità e l'originalità del suo lavoro sono stati riconosciuti a livello internazionale, come testimoniato dal conferimento del titolo di Doctor of Science honoris causa da parte della National University of Ireland nel 2006.

Il suo impegno nella comunità matematica si è espresso anche attraverso un'intensa attività istituzionale. È stato membro della Commissione Scientifica dell'Unione Matematica Italiana dal 1997 al 2015 e ha ricoperto il ruolo di Presidente della Mathesis Nazionale nel biennio 2021-2023. Inoltre, ha presieduto l'associazione AGTA (Advances in Group Theory and Applications) e ha diretto prestigiose riviste scientifiche come "Advances in Group Theory and Applications", "Ricerche di Matematica" e il "Periodico di Matematiche". Ha fatto parte dei comitati editoriali di numerose riviste internazionali, tra cui il "Mediterranean Journal of Mathematics", "Note di Matematica", "Mathematics" e il "Rendiconto dell'Accademia di Scienze Fisiche e Matematiche".

Il Professor de Giovanni ha inoltre ricoperto numerosi incarichi accademici, tra cui la Presidenza del Consiglio di Corso di Laurea in Matematica e il coordinamento

del Dottorato di Ricerca in Scienze Matematiche e Informatiche presso l'Università di Napoli Federico II. È stato membro del Consiglio Scientifico del Gruppo Nazionale per le Strutture Algebriche e Geometriche e Applicazioni dell'Istituto Nazionale di Alta Matematica (INdAM). La sua appartenenza a istituzioni di grande prestigio, come la Società Nazionale di Scienze, Lettere e Arti di Napoli e l'Accademia Peloritana dei Pericolanti, testimonia ulteriormente la sua rilevanza nella comunità scientifica.

La sua attività di divulgazione e promozione della cultura matematica ha trovato espressione nel suo ruolo di Direttore del "Periodico di Matematiche", una rivista che ha una lunga e prestigiosa storia. Fondata nel 1921 su iniziativa di Federigo Enriques e Giulio Lazzeri, il "Periodico di Matematiche" è nato come organo della Mathesis - Società Italiana di Matematica, sostituendo il precedente "Periodico di Matematica" fondato nel 1886 e cessato nel 1919. Nel corso degli anni, la rivista ha avuto alla sua guida figure di spicco come Oscar Chisini e Bruno de Finetti, i quali hanno contribuito in modo significativo alla sua valorizzazione. In questa tradizione di eccellenza, il Professor de Giovanni ha proseguito il cammino, apportando un contributo decisivo al rinnovamento e alla crescita della rivista.

Il lascito del Professor Francesco de Giovanni va ben oltre le sue pubblicazioni e i suoi incarichi istituzionali. La sua figura rappresenta un esempio di dedizione alla ricerca, all'insegnamento e alla divulgazione scientifica. La comunità matematica perde non solo uno studioso di altissimo livello, ma anche un punto di riferimento umano e professionale. La sua eredità continuerà a vivere nelle opere, nelle idee e negli insegnamenti che ha saputo trasmettere a colleghi, studenti e collaboratori.

Concludendo questo ricordo, si vuole rendere omaggio alla figura del Professor de Giovanni, il cui impegno e la cui passione per la matematica rimarranno fonte di ispirazione per le generazioni future. La sua memoria resterà viva nel cuore di tutti coloro che hanno avuto il privilegio di conoscerlo e di collaborare con lui, lasciando un segno indelebile nella storia del "Periodico di Matematiche" e della Mathesis.

Salvatore Cuomo

Direttore del Periodico di Matematiche

Domenica Di Sorbo

Presidentessa della Mathesis Nazionale

Indice di correlazione generalizzato, un'alternativa a Kendall¹

Generalized correlation index, an alternative to Kendall's

Lorenzo Millanti

Abstract

Si definisce un indice per misurare l'esistenza di un legame monotono tra due variabili, indipendentemente dalla natura del legame, sfruttando il ricorso al calcolo differenziale. L'indice che andremo a definire, infatti, usa la scala intervalare anziché soltanto la scala ordinale come fa quello di Kendall, così facendo l'indice vuole essere più informativo di quello di Kendall (il ricorso al calcolo differenziale permette di sintetizzare a tratti la tipologia del legame).

Con dati per N unità statistiche

Date $N \in \mathbb{N}^{++}$ coppie di modalità (x_i, y_i) con $i = 1, 2, \dots, N$, si sceglie di ordinare tali coppie rispetto una sola delle 2 variabili (ad es. la x) in ordine crescente, poi si calcola il differenziale intervallare di ordine 1 tra le suddette N coppie consecutive ordinate rispetto la variabile di ordinamento scelta, ottenendo così i seguenti $N-1$ differenziali:

$$Diff_i = \frac{y_{i+1} - y_i}{x_{i+1} - x_i} \text{ dove } i = 1, 2, 3, \dots, N-1.$$

Qualora ci siano denominatori pari a zero $(x_{i+1} - x_i) = 0$, allora il differenziale non verrà calcolato per quei valori dell'indice i -esimo, allo scopo di non introdurre effetti distorsivi sulla misura del legame.

Quindi, il numero di differenziali calcolabili è pari a $K \in \mathbb{N}^{++}$ dove $K \leq N - 1$ con $i = 1, 2, \dots, K$.

¹Uso del calcolo differenziale nella misurazione della correlazione generalizzata.

Si calcola poi il seguente indice di correlazione generalizzato:

$$GC = \frac{\sum_{i=1}^K \text{sgn}(Diff_i)}{N-1} = \frac{\sum_{i=1}^K \frac{|Diff_i|}{Diff_i}}{N-1}$$

Tale indice fornisce sia la presenza/assenza sia la direzione (se diretto o inverso) del legame tra le due variabili ed è un valore normalizzato compreso tra -1 e 1, vale zero in assenza di legame, vale 1 in caso di massimo legame diretto e vale -1 nel caso di massimo legame inverso.

Inoltre, poiché il valore dell'indice GC dipende della scelta della variabile di ordinamento, occorre applicare una perequazione. Ossia, dopo aver calcolato tale indice scegliendo la x come variabile di ordinamento ($=GC_{ord.x}$), si calcola nuovamente lo stesso indice ma scegliendo la y come variabile di ordinamento ($=GC_{ord.y}$), i due valori così ottenuti per GC verranno poi usati per calcolarne una media tra, il risultato così ottenuto rappresenterà il valore dell'indice GC corretto CGC.

$$CGC = \frac{GC_{ord.x} + GC_{ord.y}}{2}$$

Per approfondire le informazioni sul tipo di legame tra le due variabili (soprattutto nel caso in cui l'indice GC sia nullo per entrambe le variabili di ordinamento) potremmo ripetere gli stessi identici calcoli sopra descritti ma usando i differenziali di secondo ordine (= differenziale dei differenziali) e anche di terzo ordine, anziché limitarsi solamente al primo ordine. In questo modo l'indice di correlazione generalizzato fornisce informazioni sulla presenza di una convessità/concavità del legame (nel caso del calcolo del differenziale di secondo ordine) e sulla presenza di circolarità (nel caso del calcolo del differenziale di terzo ordine).

I differenziali di ordine superiore al primo si ottengono dai rapporti incrementali di ordine maggiore di 1, come qui di seguito riportato.

I rapporti incrementali di secondo ordine hanno come numeratore

$$\Delta^2 y_i = \Delta y_{i+1} - \Delta y_i = (y_{i+2} - y_{i+1}) - (y_{i+1} - y_i) = y_{i+2} - 2y_{i+1} + y_i$$

e come denominatore, procedendo in modo analogo al numeratore, hanno $(x_{i+2} - 2x_{i+1} + x_i)$.

I rapporti incrementali di terzo ordine hanno come numeratore

$$\Delta^3 y_i = \Delta^2 y_{i+1} - \Delta^2 y_i = y_{i+3} - 2y_{i+2} + y_{i+1} - y_{i+2} + 2y_{i+1} - y_i = y_{i+3} - 3y_{i+2} + 3y_{i+1} - y_i$$

e, analogamente, come denominatore hanno $(x_{i+3} - 3x_{i+2} + 3x_{i+1} - x_i)$.

Una seconda versione dell'indice GC che tenga in considerazione, oltre al segno, anche la magnitudine dei valori dei differenziali potrebbe essere la seguente:

$$GC_{bis} = \sum_{i=1}^K Diff_i^{1/h} \text{ dove } h \in \mathbb{N}^{++} \text{ è intero dispari e grande a piacere}$$

In tal caso l'indice varia tra $-\infty$ e $+\infty$ e vale zero in caso di assenza di legame
 Notare che:

$$\lim_{h \rightarrow +\infty} \sum_{i=1}^K Diff_i^{1/h} = \sum_{i=1}^K \lim_{h \rightarrow +\infty} Diff_i^{1/h} = K$$

Come forma normalizzata di tale indice potremmo scegliere la seguente:

$$CGC_{bis} = \operatorname{sgn}(GC_{bis}) \left[\frac{\left(1 + \frac{1}{|GC_{bis}|}\right)^{|GC_{bis}|} - 1}{e - 1} \right] = \frac{|GC_{bis}|}{GC_{bis}} \left[\frac{\left(1 + \frac{1}{|GC_{bis}|}\right)^{|GC_{bis}|} - 1}{e - 1} \right]$$

Tale indice è interpretabile in modo analogo al precedente indice di correlazione.

Anche nel caso dell'indica CGC_{bis} possiamo applicare la medesima perequazione vista in precedenza andando a calcolare i due valori dell'indice CGC_{bis} con variabile di ordinamento x ($GC_{bis_ord.x}$) e quello con variabile di ordinamento y ($GC_{bis_ord.y}$) per poi farne una media aritmetica semplice tra i due valori così ottenuti.

$$CGC_{bis} = \frac{GC_{bis_ord.x} + GC_{bis_ord.y}}{2}$$

Oltre a questo tipo di miglioramento, occorre anche gestire la differenza tra K ed $N-1$, in quanto CGC e CGC_{bis} tendono a non essere stimatori corretti nel caso in cui K tenda a 1. Qualora $K < N - 1$, si può pensare di dividere il valore di entrambi gli indici per $(N - 1 - K)$ e, nel caso particolare di CGC_{bis} , si può anche scegliere $h = 2(N - 1 - K) + 1 = 2(N - K) - 1$ per renderlo ancora più stabile e robusto e dividere lo stimatore per h .

Otterremo, quindi, i seguenti indici stabilizzati (identificati dall'asterisco).

$$GC^* = \frac{\sum_{i=1}^K \operatorname{sgn}(Diff_i)}{(N - 1) \cdot \max[(N - 1 - K), 1]} = \frac{\sum_{i=1}^K \frac{|Diff_i|}{Diff_i}}{(N - 1) \cdot \max[(N - 1 - K), 1]}$$

$$GC_{bis}^* = \frac{1}{h} \sum_{i=1}^K Diff_i^{1/h}, \quad \text{dove } h = \max[2(N - K) - 1, 1]$$

Dai quali poi derivano ovviamente anche i rispettivi valori perequati CGC^* e CGC_{bis}^* .

Con dati da tabella di contingenza

Occorre ricondurre i dati a unità statistiche in modo che tutte le coppie di dati abbiano la stessa numerosità (che per dati per unità è pari a 1) e così procedendo ci riconduciamo ai calcoli sopra esposti.

Nel caso di M variabili con $M > 2$

Occorre creare una tabella di dati per unità statistiche che sarà di dimensioni $N \times M$ dove N = numero unità statistiche e M numero di variabili esaminate, dopodiché si ordina la tabella rispetto una sola delle M variabili e rispetto quella singola variabile si calcolano i differenziali su tutte le rimanenti $M - 1$ variabili della tabella, così otterremo una matrice di differenziali totali (e non parziali) condizionatamente alla variabile di ordinamento della matrice dei dati $N \times M$ iniziale. Di ogni colonna della matrice $N - 1 \times M - 1$ dei differenziali totali così ottenuta, si calcola l'indice di correlazione generalizzato andando ad applicare il procedimento al punto 1) definito per il solo caso di due variabili e lo si applica sulle $N - 1$ righe di ciascuna delle $M - 1$ colonne; otterremo così un vettore di $M - 1$ elementi che misura il legame della variabile di ordinamento scelta rispetto tutte le altre $M - 1$ variabili. Tali $M - 1$ indici però non saranno depurati dall'effetto di interazione ed eventuali legami puri tra le M variabili.

Benchmarking degli indici.

Si esegue adesso un confronto dei valori assumibili dai vari indici in diversi contesti di dati, allo scopo di capirne il comportamento e quali siano le loro peculiarità dell'uno rispetto all'altro.

Quindi, come primo step, procediamo con il calcolo di Rho di Pearson, Tau di Kendall, CGC^* , CGC_{bis} e CGC^*_{bis} nelle tabelle sotto riportate.

Esempio 1

x	y
1	4
2	1
3	2,5
4	5
5	2,5
6	7
7	6

Rho di Pearson = 0,648675

Tau di Kendall = 0,48

$CGC^* = 0,027778$

$CGC_{bis} = 0,762497$

$CGC^*_{bis} = 0,56708$

Esempio 2

x	y
1	1
2	3
3	2
4	4
5	6
6	5
7	7
8	10
9	8
10	9

Rho di Pearson = 0,939394

Tau di Kendall = 0,82

CGC* = 0,444444

CGC_{bis} = 0,914303

CGC*_{bis} = 0,914303

Esempio 3

x	y
5	3
7	2
8	12
14	10
21	19
22	25
41	22

Rho di Pearson = 0,802083

Tau di Kendall = 0,619

CGC* = 0

CGC_{bis} = 0,140264

CGC*_{bis} = 0,140264

Esempio 4

x	y		
1	5		
1	3		
1,6	10		
2	15		
2	2	5	1
3	21	5,5	15
3	3	5,6	14,7
3,2	24	6	4
3,5	20	6	15,5
3,6	11,8	6,3	14,9
3,8	11,5	6,8	10
4	2	6,8	15,2
4	30	6,9	11
4	2	7	3
4,2	12,5	7	10,6
4,2	14	7,1	12
4,7	12,8	7,2	11,5
4,9	14,5	7,2	16
5	3	8	5

Rho di Pearson = 0,042275

Tau di Kendall = 0,045579

CGC* = -0,00714

CGC_{bis} = -0,98988

CGC*_{bis} = -0,29495

Esempio 5

x	y		
5,01	1		
5,01	3		
5,01	5		
5,01	7		
5,01	9	5,02	6
5,01	11	5,02	8
5,01	13	5,02	10
5,01	15	5,02	12
5,01	17	5,02	14
5,01	19	5,02	16
5,01	21	5,02	18
5,01	23	5,02	20
5,01	25	5,02	22
5,01	27	5,02	24
5,01	29	5,02	26
5,01	31	5,02	28
5,01	33	5,02	30
5,02	2	5,02	32
5,02	4	5,02	34

Rho di Pearson = 0,050965

Tau di Kendall = 0,04222

CGC* = 0,014918

CGC_{bis} = -0,48613

CGC*_{bis} = -0,00761

Esempio 6

x	y
1	0,7
2	0,4
3	1
4	1
5	1,1
6	1,2
8	1,2
10	1,2

Rho di Pearson = 0,781809

Tau di Kendall = 0,848668

CGC* = 0,163265

CGC_{bis} = 0,696708

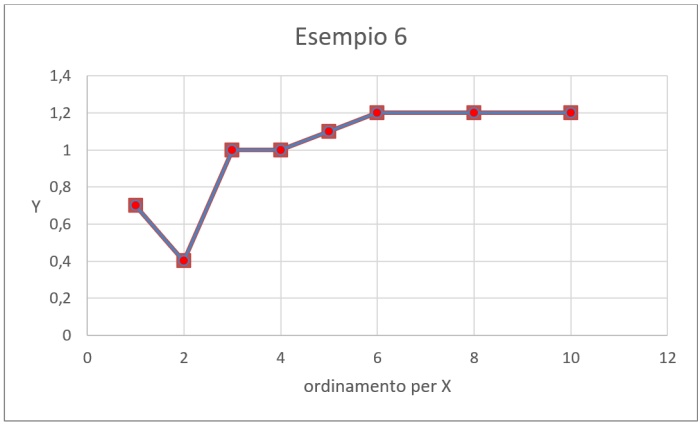
CGC*_{bis} = 0,213018

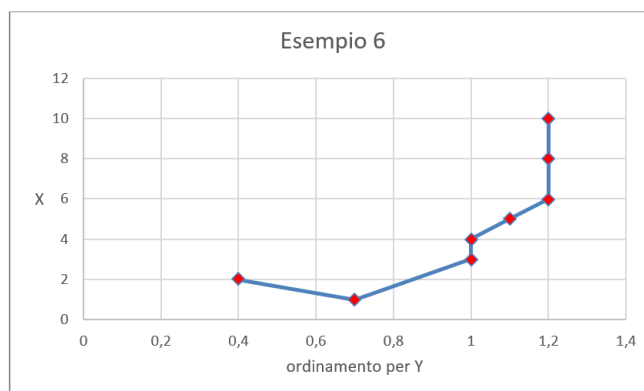
Successivamente, si passa al secondo step della fase analitica del confronto, ricorrendo alla seguente analisi discriminante dei risultati delle precedenti stime.

Facendo sintesi dei risultati ottenuti osserviamo che, rispetto un prudenziale cut-off in valore assoluto di 0,2 per gli indici normalizzati, esistono allineamenti e disallineamenti tra l'individuazione della presenza/assenza dei legami, come riportato dalla tabella sottostante.

		presenza legame per CGC*bis	
		si	no (<0,2)
presenza legame per Rho e/o Tau	si	esempi 1 e 2	esempio 3
	no (<0,2)	esempio 4	esempio 5

Degno di particolare nota è l'esempio 6, in quanto di difficile collocazione all'interno della precedente tabella di confronto; tale esempio fornisce la casistica in cui sia Rho che Tau sono elevati mentre CGC_{bis} è elevato e CGC_{bis}^* è prossimo a zero. Da un'osservazione grafica dei dati dell'esempio 6, un ramo della spezzata che unisce i punti presenta un certo andamento monotono ma un altro ramo presenta totale assenza di legame (sia ordinando per X che per Y), quindi possiamo constatare che CGC è molto sensibile a rilevare ambiguità dal punto di vista della monotonia dei dati, informazioni che dall'uso di Rho e di Tau non è possibile evincere.





Di conseguenza, in caso di valori di segno concorde ma sensibilmente diversi tra CGC_{bis} e CGC^{*}_{bis} ,

possiamo ipotizzare la presenza di legami non così netti e certi come invece Rho e Tau rilevano, in questo caso, potremmo anche introdurre un'ulteriore perequazione del nostro nuovo indice data dalla media aritmetica semplice tra CGC_{bis} e CGC^{*}_{bis} .

L'esempio 4, invece, mette in evidenza la capacità dell'indice CGC_{bis} o CGC^{*}_{bis} di percepire la presenza di un andamento parabolico dei dati, cosa che gli altri due indici Rho e Tau non rilevano.

L'esempio 3 rileva una monotonia a segni alterni e questo rende infinitesimo CGC_{bis} e CGC^{*}_{bis} .

In generale si osserva che CGC_{bis} e CGC^{*}_{bis} perdono attendibilità in caso di andamenti oscillatori in cui la monotonia ha un segno alternato al variare di i da 1 a $N-1$, risultano inoltre molto sensibili alla presenza di rami delle serie di dati con andamento approssimativamente parallelo all'asse delle ascisse, ma sono altresì capaci (proprio in virtù della loro sensibilità verso la presenza/assenza della monotonia) di individuare andamenti simili a quelli delle coniche (vedi ad es. la parabola con concavità rivolta verso il basso dell'esempio 4). Quindi, in estrema sintesi, rispetto Rho e Tau, i due indici CGC_{bis} e CGC^{*}_{bis} mostrano una maggiore sensibilità verso la presenza/assenza di comportamenti monotoni dei dati.

Uno degli esempi nella realtà in cui tale indice potrebbe funzionare meglio degli altri è nella medicina, qualora si voglia fare uno studio esplorativo per individuare un qualche legame tra le misure ottenute da due esami diagnostici eseguiti su un campione di individui. Il precedente esempio 4 è, infatti, tratto dalla misurazione di due valori del sangue: x = Calcitonina, y = Creatinichinasi.

Circocentro e Ortocentro di un triangolo

Circumcenter and Orthocenter of a Triangle

Franco Fernicola¹

Abstract

Le formule che esplicherò in quest'articolo riguardano due punti notevoli di un triangolo, ovvero il *circocentro* e l'*ortocentro*. Su molti testi si trovano relazioni sulle coordinate del *baricentro* e dell'*incentro* in funzione delle coordinate dei vertici, non trovando relazioni analoghe in circolazione su *circocentro* e *ortocentro* ho ritenuto interessante indagare in tale direzione. Per le coordinate del *circocentro* si intersecano due assi di due lati qualunque, allo stesso modo per l'*ortocentro* determiniamo le coordinate del punto di intersezione di due rette che contengono due qualsiasi altezze. Le formule sono interessanti per la "loro simmetria" se esplicitate mediante i determinanti di matrici di ordine 3.

Introduzione

Quando si parla di un triangolo, che risulta essere il poligono con il minimo numeri di lati, è inevitabile non prendere in considerazione alcuni suoi punti notevoli che sono noti fin dall'antichità ovvero *baricentro*, *incentro*, *circocentro*, *ortocentro* ed *excentro*. Questi citati sono forse i più noti e vengono evidenziati anche nel corso degli studi della scuola secondaria superiore. Esiste in realtà una lista di oltre 10.000 punti associabili ad un triangolo e consultabile nell'*Encyclopedia of Triangle Centers* curata dal matematico *Clark Kimberling*. Con l'introduzione del metodo delle coordinate i vertici di un triangolo sono individuati da coppie ordinate di numeri reali e alcuni punti notevoli in precedenza citati sono descritti dalle coordinate dei vertici (baricentro), dalle coordinate dei vertici e dalle lunghezze dei lati (incentro). Per determinare le coordinate degli altri punti notevoli quali circocentro, ortocentro ed excentro in funzione dei vertici è necessario applicare le definizioni e svolgere i vari calcoli. Questo articolo nasce con l'intento di esplicitare le coordinate di *circocentro* e *ortocentro* in funzione delle coordinate dei vertici assegnati. Così come assegnate le coordinate di

¹Liceo Assteas - Buccino (SA) - franco.fernicola@iisassteas.edu.it

tre punti non allineati $A \equiv (x_1, y_1)$, $B \equiv (x_2, y_2)$ e $C \equiv (x_3, y_3)$ è possibile determinare le coordinate dell'*incentro* e del *baricentro*, in modo analogo dimostreremo che è possibile determinare le coordinate del *circocentro* $H \equiv (H_x, H_y)$ e dell'*ortocentro* $K \equiv (K_x, K_y)$ del triangolo in funzione dei vertici assegnati. Nei calcoli si utilizza il metodo di Cramer per la risoluzione di un sistema lineare, le stesse relazioni trovate possono essere scritte mediante determinanti di matrici del 3° ordine per apprezzarne la simmetria nella scrittura. Per completezza anteponiamo anche le formule per l'*incentro* e il *baricentro* in funzione dei vertici.

Incentro di un triangolo

Definizione 1. Si chiama *incentro* di un triangolo il punto di incontro di due qualsiasi bisettrici. Assegnate le coordinate dei tre vertici $A \equiv (x_1, y_1)$, $B \equiv (x_2, y_2)$ e $C \equiv (x_3, y_3)$ è noto che le coordinate dell'*incentro* I sono date dalla seguente relazione:

$$I \equiv \left(\frac{ax_1 + bx_2 + cx_3}{2p}, \frac{ay_1 + by_2 + cy_3}{2p} \right)$$

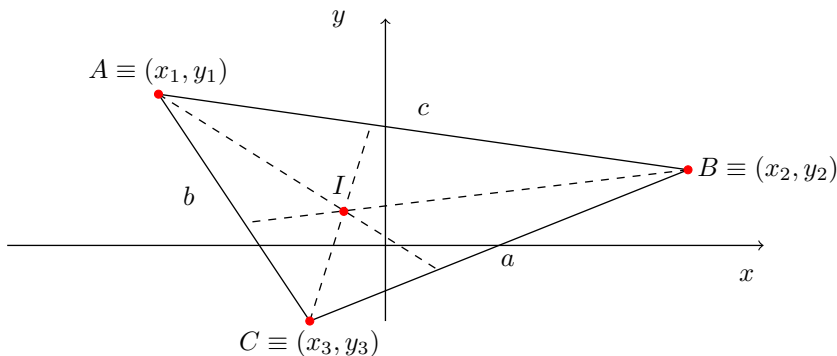
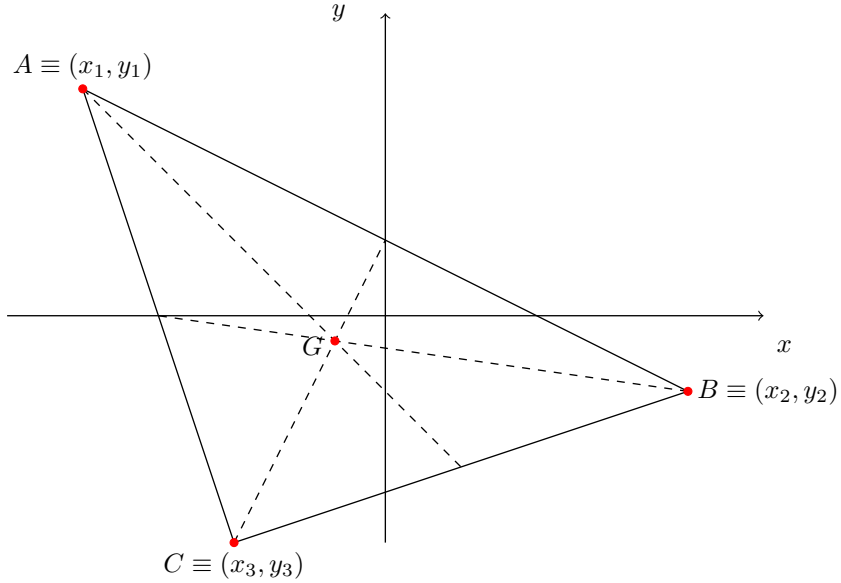


Figura 1: Incentro I del triangolo

Baricentro di un triangolo

Definizione 2. Si chiama *baricentro* di un triangolo il punto di incontro di due qualsiasi mediane. Assegnate le coordinate dei tre vertici $A \equiv (x_1, y_1)$, $B \equiv (x_2, y_2)$ e $C \equiv (x_3, y_3)$ è noto che le coordinate del baricentro G sono date dalla seguente relazione:

$$G \equiv \left(\frac{x_1 + x_2 + x_3}{3}, \frac{y_1 + y_2 + y_3}{3} \right)$$

Figura 2: Baricentro G del triangolo

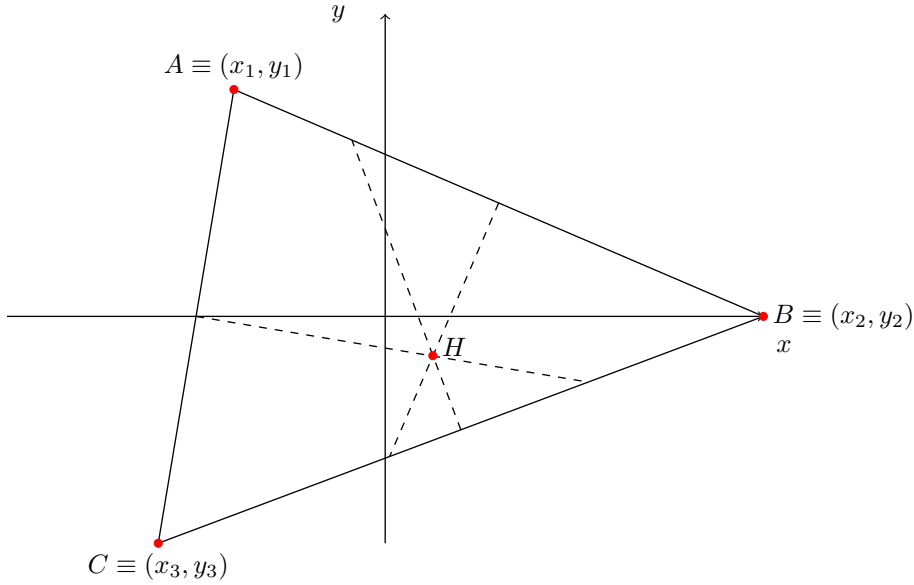
Osservazione 1. *Precisiamo che in seguito utilizzeremo le lettere a, b, c, d, e, f per denotare le coordinate dei vertici e rendere la scrittura più snella nel dimostrare le formule, l'uso degli indici x_i e y_i con $1 \leq i \leq 3$ appesantirebbe di molto la scrittura delle relazioni, cosa che invece faremo alla fine di ciascuna sezione scrivendo le formule determinate con una notazione che dipende dalle coordinate dei vertici: $A \equiv (x_1, y_1)$, $B \equiv (x_2, y_2)$, $C \equiv (x_3, y_3)$ così da farne apprezzare "una certa simmetria".*

Circocentro di un triangolo

Definizione 3. *Si chiama circocentro del triangolo il punto di incontro dei tre assi. Ci chiediamo se conoscendo le coordinate dei tre vertici $A \equiv (a, b)$, $B \equiv (b, c)$ e $C \equiv (d, e)$, possiamo conoscere le coordinate del circocentro $H = (H_x, H_y)$ del triangolo? La risposta è affermativa e le coordinate del circocentro sono date dalle seguenti relazioni:*

$$H_x = \frac{(a^2 + b^2)(d - f) - (c^2 + d^2)(b - f) + (e^2 + f^2)(b - d)}{2[a(d - f) - c(b - f) + e(b - d)]}$$

$$H_y = \frac{(a^2 + b^2)(e - c) - (c^2 + d^2)(e - a) + (e^2 + f^2)(c - a)}{2[a(d - f) - c(b - f) + e(b - d)]}$$

Figura 3: Circocentro H del triangolo

Dimostrazione. Sappiamo che la retta che passa per AB ha equazione $(b-d)x + (c-a)y + ad - bc = 0$. La retta ortogonale a questa passante per il punto medio $M \equiv \left(\frac{a+c}{2}, \frac{b+d}{2}\right)$ del segmento AB ha equazione $(a-c)x + (b-d)y = \frac{a^2-c^2}{2} + \frac{b^2-d^2}{2}$. Allo stesso modo la retta che passa per AC ha equazione $(b-f)x + (e-a)y + af - be = 0$. La retta ortogonale a questa passante per il punto medio $N \equiv \left(\frac{a+e}{2}, \frac{b+f}{2}\right)$ del segmento AC ha equazione $(a-e)x + (b-f)y = \frac{a^2-e^2}{2} + \frac{b^2-f^2}{2}$. In corrispondenza del seguente sistema lineare si trova il circo-

$$\text{centro: } \begin{cases} (a-c)x + (b-d)y = \frac{a^2-c^2}{2} + \frac{b^2-d^2}{2} \\ (a-e)x + (b-f)y = \frac{a^2-e^2}{2} + \frac{b^2-f^2}{2} \end{cases}.$$

$$\Delta = \begin{vmatrix} (a-c) & (b-d) \\ (a-e) & (b-f) \end{vmatrix} = a(d-f) - c(b-f) + e(b-d)$$

$$\Delta_x = \begin{vmatrix} \frac{a^2-c^2}{2} + \frac{b^2-d^2}{2} & b-d \\ \frac{a^2-e^2}{2} + \frac{b^2-f^2}{2} & b-f \end{vmatrix} = \begin{vmatrix} \frac{a^2+b^2}{2} - \frac{c^2+d^2}{2} & b-d \\ \frac{a^2+b^2}{2} - \frac{e^2+f^2}{2} & b-f \end{vmatrix} =$$

$$= \frac{1}{2}(a^2 + b^2) \begin{vmatrix} 1 & b-d \\ 1 & b-f \end{vmatrix} - \frac{1}{2} \begin{vmatrix} c^2 + d^2 & b-d \\ e^2 + f^2 & b-f \end{vmatrix} =$$

$$\frac{1}{2}(a^2 + b^2)(d-f) - \frac{1}{2}(c^2 + d^2)(b-f) + \frac{1}{2}(e^2 + f^2)(b-d).$$

$$\Delta_y = \begin{vmatrix} a-c & \frac{a^2-c^2}{2} + \frac{b^2-d^2}{2} \\ a-e & \frac{a^2-e^2}{2} + \frac{b^2-f^2}{2} \end{vmatrix} = \begin{vmatrix} a-c & \frac{a^2+b^2}{2} - \frac{c^2+d^2}{2} \\ a-e & \frac{a^2+b^2}{2} - \frac{e^2+f^2}{2} \end{vmatrix} =$$

$$= \frac{1}{2}(a^2 + b^2) \begin{vmatrix} a-c & 1 \\ a-e & 1 \end{vmatrix} - \frac{1}{2} \begin{vmatrix} a-c & c^2+d^2 \\ a-e & e^2+f^2 \end{vmatrix} =$$

$$= \frac{1}{2}(a^2 + b^2)(e-c) - \frac{1}{2}(c^2 + d^2)(e-a) + \frac{1}{2}(e^2 + f^2)(c-a).$$

Possiamo concludere:

$$H_x = \frac{\Delta_x}{\Delta} = \frac{(a^2 + b^2)(d-f) - (c^2 + d^2)(b-f) + (e^2 + f^2)(b-d)}{2[a(d-f) - c(b-f) + e(b-d)]} = \frac{1}{2} \begin{vmatrix} a^2+b^2 & b & 1 \\ c^2+d^2 & d & 1 \\ e^2+f^2 & f & 1 \end{vmatrix} \begin{vmatrix} a & b & 1 \\ c & d & 1 \\ e & f & 1 \end{vmatrix}$$

$$H_y = \frac{\Delta_y}{\Delta} = \frac{(a^2 + b^2)(e-c) - (c^2 + d^2)(e-a) + (e^2 + f^2)(c-a)}{2[a(d-f) - c(b-f) + e(b-d)]} = -\frac{1}{2} \begin{vmatrix} a^2+b^2 & a & 1 \\ c^2+d^2 & c & 1 \\ e^2+f^2 & e & 1 \end{vmatrix} \begin{vmatrix} a & b & 1 \\ c & d & 1 \\ e & f & 1 \end{vmatrix}$$

□

Osserviamo che i tre punti essendo non allineati il determinante a denominatore è non nullo. Se per i punti utilizziamo la notazione $A \equiv (x_1, y_1)$, $B \equiv (x_2, y_2)$ e $C \equiv (x_3, y_3)$, abbiamo le seguenti relazioni:

$$H_x = \frac{1}{2} \frac{\begin{vmatrix} x_1^2 + y_1^2 & y_1 & 1 \\ x_2^2 + y_2^2 & y_2 & 1 \\ x_3^2 + y_3^2 & y_3 & 1 \end{vmatrix}}{\begin{vmatrix} x_1 & y_1 & 1 \\ x_2 & y_2 & 1 \\ x_3 & y_3 & 1 \end{vmatrix}} \quad H_y = -\frac{1}{2} \frac{\begin{vmatrix} x_1^2 + y_1^2 & x_1 & 1 \\ x_2^2 + y_2^2 & x_2 & 1 \\ x_3^2 + y_3^2 & x_3 & 1 \end{vmatrix}}{\begin{vmatrix} x_1 & y_1 & 1 \\ x_2 & y_2 & 1 \\ x_3 & y_3 & 1 \end{vmatrix}}$$

Ortocentro di un triangolo

Definizione 4. Si chiama ortocentro del triangolo il punto di incontro delle tre altezze. Ci chiediamo se conoscendo le coordinate dei tre vertici $A \equiv (a, b)$, $B \equiv (b, c)$ e $C \equiv (d, e)$, possiamo conoscere le coordinate dell'ortocentro $K = (K_x, K_y)$ del triangolo? La risposta è affermativa e le coordinate dell'ortocentro sono date dalle seguenti relazioni:

$$K_x = \frac{\Delta_x}{\Delta} = \frac{(b^2 + ce)(f - d) - (d^2 + ae)(f - b) + (f^2 + ac)(d - b)}{a(d - f) - c(b - f) + e(b - d)}$$

$$K_y = \frac{\Delta_y}{\Delta} = \frac{(a^2 + df)(c - e) - (c^2 + bf)(a - e) + (e^2 + bd)(a - c)}{a(d - f) - c(b - f) + e(b - d)}$$

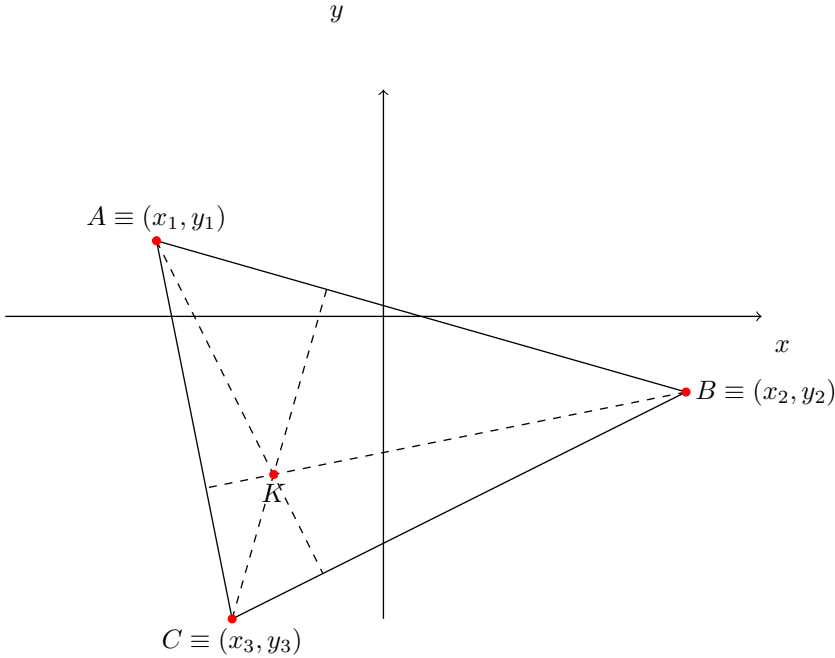


Figura 4: Ortocentro K del triangolo

Dimostrazione. Sappiamo che la retta che passa per AB ha equazione $(b - d)x + (c - a)y + ad - bc = 0$. La retta ortogonale a questa passante per il punto $C \equiv (e, f)$ ha equazione $(a - c)x + (b - d)y = (a - c)e + (b - d)f$. Allo stesso modo la retta che passa per AC ha equazione $(b - f)x + (e - a)y + af - be = 0$. La retta ortogonale a questa passante per il punto medio $B \equiv (c, d)$ ha equazione

$$(a - e)x + (b - f)y = (a - e)c + (b - f)d.$$

In corrispondenza del seguente sistema lineare si trova l'ortocentro:

$$\begin{cases} (a - c)x + (b - d)y = (a - c)e + (b - d)f \\ (a - e)x + (b - f)y = (a - e)c + (b - f)d \end{cases}.$$

$$\Delta = \begin{vmatrix} a - c & b - d \\ a - e & b - f \end{vmatrix} = a(b - f) - c(b - f) - a(b - d) + e(b - d) = a(d - f) - c(b - f) + e(b - d)$$

$$\Delta_x = \begin{vmatrix} (a - c)e + (b - d)f & b - d \\ (a - e)c + (b - f)d & b - f \end{vmatrix} = \begin{vmatrix} (a - c)e & b - d \\ (a - e)c & b - f \end{vmatrix} + \begin{vmatrix} (b - d)f & b - d \\ (b - f)d & b - f \end{vmatrix} =$$

$$= \begin{vmatrix} ae - ce & b - d \\ ac - ce & b - f \end{vmatrix} + (b - d)(b - f) \begin{vmatrix} f & 1 \\ d & 1 \end{vmatrix} = \begin{vmatrix} ae & b - d \\ ac & b - f \end{vmatrix} - \begin{vmatrix} ce & b - d \\ ce & b - f \end{vmatrix} + (b - d)(b - f) \begin{vmatrix} f & 1 \\ d & 1 \end{vmatrix} =$$

$$= \begin{vmatrix} ae & b - d \\ ac & b - f \end{vmatrix} - ce \begin{vmatrix} 1 & b - d \\ 1 & b - f \end{vmatrix} + (b - d)(b - f)(f - d) =$$

$$= ae(b - f) - ac(b - d) + ce(f - d) + (b - d)(b - f)(f - d) =$$

$$= ae(b - f) - ac(b - d) + ce(f - d) + b^2(f - d) + d^2(b - f) + f^2(d - b) =$$

$$= (b^2 + ce)(f - d) - (d^2 + ae)(f - b) + (f^2 + ac)(d - b)$$

$$\Delta_y = \begin{vmatrix} a - c & (a - c)e + (b - d)f \\ a - e & (a - e)c + (b - f)d \end{vmatrix} = \begin{vmatrix} a - c & (a - c)e \\ a - e & (a - e)c \end{vmatrix} + \begin{vmatrix} a - c & (b - d)f \\ a - e & (b - f)d \end{vmatrix} =$$

$$= (a - c)(a - e) \begin{vmatrix} 1 & e \\ 1 & c \end{vmatrix} + \begin{vmatrix} a - c & bf - df \\ a - e & bd - df \end{vmatrix} =$$

$$= (a - c)(a - e) \begin{vmatrix} 1 & e \\ 1 & c \end{vmatrix} + \begin{vmatrix} a - c & bf \\ a - e & bd \end{vmatrix} - \begin{vmatrix} a - c & df \\ a - e & df \end{vmatrix} =$$

$$= (a - c)(a - e) \begin{vmatrix} 1 & e \\ 1 & c \end{vmatrix} + \begin{vmatrix} a - c & bf \\ a - e & bd \end{vmatrix} - df \begin{vmatrix} a - c & 1 \\ a - e & 1 \end{vmatrix} =$$

$$\begin{aligned}
&= (a-c)(a-e)(c-e) + bd(a-c) - bf(a-e) + df(c-e) = \\
&= a^2(c-e) - c^2(a-e) + e^2(a-c) + bd(a-c) - bf(a-e) + df(c-e) = \\
&= (a^2 + df)(c-e) - (c^2 + bf)(a-e) + (e^2 + bd)(a-c).
\end{aligned}$$

Possiamo concludere:

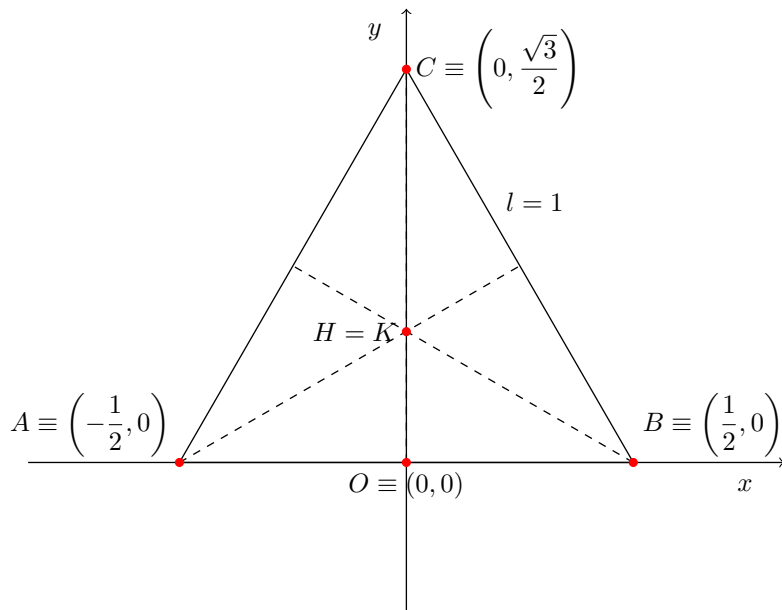
$$\begin{aligned}
K_x = \frac{\Delta_x}{\Delta} &= \frac{(b^2 + ce)(f-d) - (d^2 + ae)(f-b) + (f^2 + ac)(d-b)}{a(d-f) - c(b-f) + e(b-d)} = - \frac{\begin{vmatrix} b^2 + ce & b & 1 \\ d^2 + ae & d & 1 \\ f^2 + ac & f & 1 \end{vmatrix}}{\begin{vmatrix} a & b & 1 \\ c & d & 1 \\ e & f & 1 \end{vmatrix}} \\
K_y = \frac{\Delta_y}{\Delta} &= \frac{(a^2 + df)(c-e) - (c^2 + bf)(a-e) + (e^2 + bd)(a-c)}{a(d-f) - c(b-f) + e(b-d)} = \frac{\begin{vmatrix} a^2 + df & a & 1 \\ c^2 + bf & c & 1 \\ e^2 + bd & e & 1 \end{vmatrix}}{\begin{vmatrix} a & b & 1 \\ c & d & 1 \\ e & f & 1 \end{vmatrix}}
\end{aligned}$$

Osserviamo che i tre punti essendo non allineati il determinante a denominatore è non nullo. Se per i punti utilizziamo la notazione $A \equiv (x_1, y_1)$, $B \equiv (x_2, y_2)$ e $C \equiv (x_3, y_3)$, abbiamo le seguenti relazioni:

$$\begin{aligned}
K_x &= - \frac{\begin{vmatrix} y_1^2 + x_2x_3 & y_1 & 1 \\ y_2^2 + x_1x_3 & y_2 & 1 \\ y_3^2 + x_1x_2 & y_3 & 1 \end{vmatrix}}{\begin{vmatrix} x_1 & y_1 & 1 \\ x_2 & y_2 & 1 \\ x_3 & y_3 & 1 \end{vmatrix}} & K_y &= \frac{\begin{vmatrix} x_1^2 + y_2y_3 & x_1 & 1 \\ x_2^2 + y_1y_3 & x_2 & 1 \\ x_3^2 + y_1y_2 & x_3 & 1 \end{vmatrix}}{\begin{vmatrix} x_1 & y_1 & 1 \\ x_2 & y_2 & 1 \\ x_3 & y_3 & 1 \end{vmatrix}}
\end{aligned}$$

□

Esempio 1. Supponiamo che in un sistema di riferimento abbiamo tre punti con le seguenti coordinate: $A \equiv \left(-\frac{1}{2}, 0\right)$, $B \equiv \left(\frac{1}{2}, 0\right)$ e $C \equiv \left(0, \frac{\sqrt{3}}{2}\right)$. Abbiamo un triangolo equilatero avente lato $l = 1$ e altezza $h = \frac{\sqrt{3}}{2}$. Ovviamente circocentro e ortocentro coincidono, per ovvie considerazioni geometriche il segmento $\overline{OK}$ è la terza parte dell'altezza e di conseguenza $H = K \equiv \left(0, \frac{\sqrt{3}}{6}\right)$.

Figura 5: Ortocentro e Circocentro $H = K$ del triangolo

$$H_x = \frac{1}{2} \begin{vmatrix} \frac{1}{4} & 0 & 1 \\ \frac{1}{4} & 0 & 1 \\ \frac{3}{4} & \frac{\sqrt{3}}{2} & 1 \\ -\frac{1}{2} & 0 & 1 \\ \frac{1}{2} & 0 & 1 \\ 0 & \frac{\sqrt{3}}{2} & 1 \end{vmatrix} = 0$$

$$H_y = -\frac{1}{2} \begin{vmatrix} \frac{1}{4} & -\frac{1}{2} & 1 \\ \frac{1}{4} & \frac{1}{2} & 1 \\ \frac{3}{4} & 0 & 1 \\ -\frac{1}{2} & 0 & 1 \\ \frac{1}{2} & 0 & 1 \\ 0 & \frac{\sqrt{3}}{2} & 1 \end{vmatrix} = \left(-\frac{1}{2}\right).$$

$$\frac{\frac{1}{8} - \frac{3}{8} - \frac{3}{8} + \frac{1}{8}}{\frac{\sqrt{3}}{2}} = \frac{\sqrt{3}}{6}$$

$$K_x = - \frac{\begin{vmatrix} 0 & 0 & 1 \\ 0 & 0 & 1 \\ \frac{1}{2} & \frac{\sqrt{3}}{2} & 1 \\ -\frac{1}{2} & 0 & 1 \\ \frac{1}{2} & 0 & 1 \\ 0 & \frac{\sqrt{3}}{2} & 1 \end{vmatrix}}{\begin{vmatrix} 0 & 0 & 1 \\ 0 & 0 & 1 \\ -\frac{1}{2} & 0 & 1 \\ \frac{1}{2} & 0 & 1 \\ 0 & \frac{\sqrt{3}}{2} & 1 \end{vmatrix}} = 0 \quad K_y = \frac{\begin{vmatrix} \frac{1}{4} & -\frac{1}{2} & 1 \\ \frac{1}{4} & \frac{1}{2} & 1 \\ 0 & 0 & 1 \\ -\frac{1}{2} & 0 & 1 \\ \frac{1}{2} & 0 & 1 \\ 0 & \frac{\sqrt{3}}{2} & 1 \end{vmatrix}}{\begin{vmatrix} 0 & 0 & 1 \\ 0 & 0 & 1 \\ -\frac{1}{2} & 0 & 1 \\ \frac{1}{2} & 0 & 1 \\ 0 & \frac{\sqrt{3}}{2} & 1 \end{vmatrix}} = \frac{\frac{1}{4}}{\frac{\sqrt{3}}{2}} = \frac{\sqrt{3}}{6}$$

Bibliografia

- [1] Emilio Ambrisi. Geometria euclidea dopo euclide. Periodico di Matematiche, (2), 2004.
- [2] D Besso. Di una serie di punti notevoli nel triangolo. Periodico di Matematiche, 2(1), 2018.
- [3] Marcello Bruni. Un problema sui punti notevoli del triangolo. Periodico di Matematiche, 3(3), 1996.
- [4] Lucio Cadeddu and Maria Pia Gusi. Napoleone, fermat, e i problemi dei triangoli. Lettera Matematica Pristem, 96:26–30, 2016.
- [5] Alessandra Fiocca. Il problema di malfatti nella letteratura matematica dell’800. Annali dell’Universit’a di Ferrara, 26:173–202, 1980.
- [6] Cesare Labianca. I punti isodinamici di un triangolo. PERIODICO DI MATEMATICA, page 147, 2022.
- [7] Elena Lazzari. Laboratorio di geometria origami: prime considerazioni su un’indagine esplorativa. Annali online della Didattica e della Formazione Docente, 14(24):123–143, 2022.
- [8] Lorenzo Mazza et al. Triangoli, circonferenze e punti notevoli. PROGETTO ALICE, 14:113–130, 2013.
- [9] Loredana Rossi. I luoghi geometrici attraverso le costruzioni. EUT Edizioni Universit’a di Trieste, 2011.
- [10] Libero Verardi et al. La retta d’eulero. IPOTESI, 7:7–10, 2004.

Crittografia: una Prospettiva Didattica – parte 1

Cryptography: an Educational Perspective – part 1

Nicola Fusco¹

Abstract

Cryptography is the ancient discipline that deals with developing methods for transforming text messages so that they are incomprehensible to anyone other than the legitimate recipient. In this work we will look at a small selection of classical cryptographic techniques that have links with high school mathematics topics. Given the importance that cryptography has assumed in the ubiquitous digital communications, this topic can constitute a module of civics education.

Introduzione

Ogni tecnica crittografica consiste nell'operazione di *cifratura*, con cui un messaggio in chiaro viene trasformato in un messaggio incomprensibile, e in quella inversa di *decifratura*.

La *crittografia a chiave privata* è caratterizzata dal fatto che la cifratura e la decifratura richiedono una stessa informazione (la *chiave*), da usare sia in fase di cifratura sia in fase di decifratura: nella seconda si eseguono le stesse operazioni della prima ma in ordine inverso, per cui questa forma è detta anche *crittografia simmetrica*. L'acquisizione di tale chiave permette l'immediata decodifica del messaggio da parte di chiunque e quindi tale chiave deve essere scambiata tra i due soggetti in modo molto sicuro.²

¹Docente di Matematica e Fisica presso il Liceo Scientifico Statale "A. Scacchi" Bari (BA), fusco.nicola@liceoscacchibari.it.

²Una conseguenza di ciò è che, nella crittografia a chiave privata, le chiavi devono essere stabilite e fornite in anticipo ai diversi "interlocutori" della conversazione cifrata. Finché le chiavi erano molto semplici esse potevano venire spedite con canali diversi rispetto ai messaggi cifrati, pur con il rischio che venissero intercettate. Oppure venivano consegnate in particolari occasioni in cui gli interlocutori si trovavano nello stesso luogo. Tuttavia ciò implica il non poter cambiare chiavi in modo estemporaneo e anche la necessità di usare la stessa chiave per un tempo più o meno lungo.

Nel seguito sono esposti tre metodi di cifratura a chiave privata (con griglia, per sostituzione con parola chiave e mediante matrici). In tutti i casi sono evidenziati i collegamenti con argomenti di matematica svolti nelle scuole secondarie di II grado. Le appendici sono disponibili presso l'autore.

1. Cifratura con griglia

Questo sistema, a chiave privata, fa parte dei metodi di *cifratura per trasposizione*: le lettere del messaggio in chiaro e quelle del messaggio cifrato sono le stesse, ma cambiate di posto in modo da essere incomprensibili.

In questo metodo si costruisce una griglia quadrata $n \times n$ in modo che il numero di celle sia maggiore o uguale al numero di lettere del messaggio da cifrare. È necessario che n sia pari affinché la griglia sia suddivisibile in quattro quadranti uguali. Per un uso migliore del metodo è opportuno costruire questa griglia su un supporto di cartone con una cornice intorno contrassegnando in qualche modo uno degli angoli.

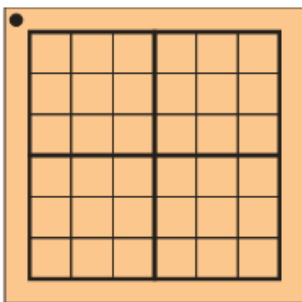


Figura 1: Figura Griglia quadrata per cifratura : sono evidenziati i quattro quadranti (ma non è necessario che siano evidenziati anche nella realizzazione della griglia).

Dobbiamo ora creare uno schema di $\frac{n^2}{4}$ fori su questa griglia per mescolare le lettere del messaggio. Questo mescolamento avviene ruotando la griglia di multipli interi di un angolo retto e scrivendo le lettere del messaggio nelle posizioni dei fori dopo ogni rotazione. Di conseguenza lo schema dei fori deve essere tale che nessuna delle quattro posizioni che un dato foro assume con le varie rotazioni coincida con la posizione di un altro foro in una delle altre rotazioni. Per ottenere questo risultato scegliamo una casella qualunque, la ritagliamo via e poi contrassegniamo con una X le caselle che corrispondono, per rotazione, a quella scelta negli altri quadranti. Ad esempio

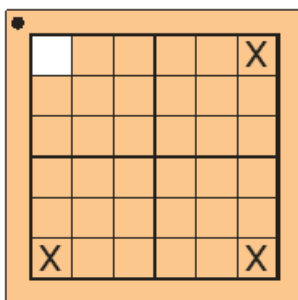


Figura 2: Figura Griglia con il primo foro e l'indicazione delle quattro caselle corrispondenti per rotazione, da non forare nelle operazioni successive.

Ora scegliamo un'altra casella e facciamo la stessa cosa.

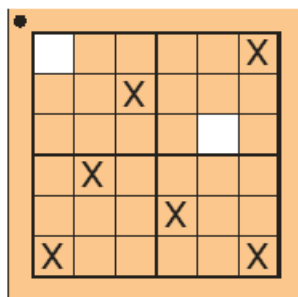


Figura 3: Griglia con i primi due fori e le relative caselle da non forare.

Proseguiamo così fino a forare o contrassegnare tutte le caselle della griglia.

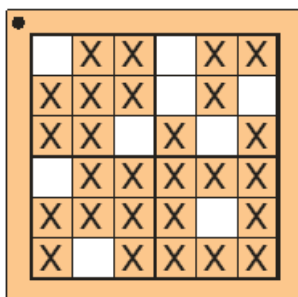


Figura 4: Griglia di cifratura completa.

A questo punto siamo pronti a cifrare il nostro messaggio: "Frequento il Liceo Scientifico A. SCACCHI". La prima cosa è scrivere la frase senza spazi, punteggiatura, accenti e con tutte le lettere maiuscole o tutte minuscole:³

FREQUENTOILLICEOSCIENTIFICOASCACCHI

Poi posiamo la griglia su un foglio di carta, segniamo sul foglio i quattro angoli (in modo da avere dei riferimenti per le successive rotazioni), e poi cominciamo a riempire i fori con le lettere della frase.

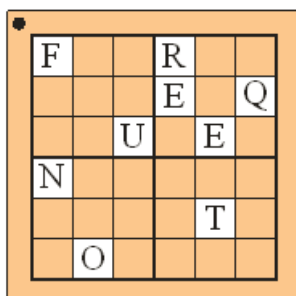


Figura 5: Figura Posizione iniziale della griglia durante la cifratura.

A questo punto ruotiamo la griglia di un angolo retto in senso orario e poi continuiamo a scrivere la frase, proseguendo con le successive rotazioni e scritture. Se la frase è più corta delle caselle a disposizione, riempiamo le ultime con delle lettere a caso.

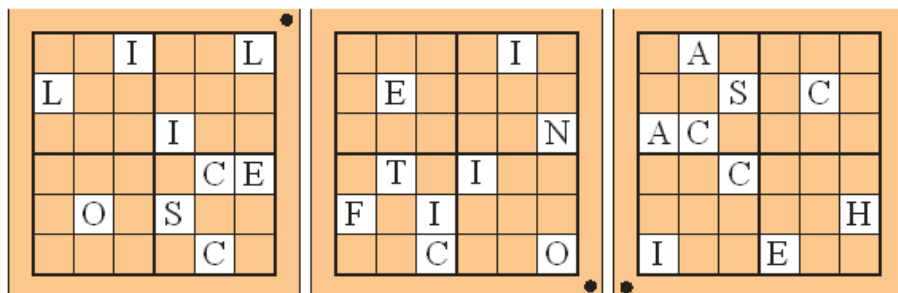


Figura 6: Seconda, terza e quarta posizione della griglia durante la cifratura.

Dopo aver completato questa operazione, rimuovendo la griglia, ci si trova davanti ad un quadrato di lettere in cui sono mescolate, apparentemente a caso, le lettere del messaggio originale.

³Questa procedura serve a rendere più difficile la decrittazione da parte di destinatari non designati.

F	A	I	R	I	L
L	E	S	E	C	Q
A	C	U	I	E	N
N	T	C	I	C	E
F	O	I	S	T	H
I	O	C	E	C	O

Figura 7: Figura Messaggio cifrato con griglia.

Copiando secondo il normale verso di lettura si ha il messaggio cifrato:

FAIRILLESECQACUIENNTCICEFOISTHIOCECO

Chi riceve il messaggio deve avere necessariamente una copia della griglia e, una volta ricevuto il messaggio, lo può decodificare procedendo al contrario:

1. Scrive il messaggio in una griglia quadrata,
2. Dispone la griglia nella posizione iniziale, legge le lettere e le scrive in sequenza da sinistra a destra e dall'alto in basso,
3. Ruota la griglia di un angolo retto in senso orario,
4. Ripete la sequenza 2)-3) fino a completare tutte le rotazioni e rilevare tutte le lettere nel mondo giusto.

1.1. Vantaggi e svantaggi di questo metodo

Questo metodo è molto veloce, sia nella cifratura sia nella decifratura, tuttavia ciò si paga con una vulnerabilità maggiore. In particolare tale vulnerabilità risiede nel fatto che, come in tutti i metodi di cifratura per trasposizione, il messaggio cifrato presenta le stesse identiche lettere di quello in chiaro, per cui con tecniche combinatorie si può riuscire a decifrare il messaggio, anche ragionando a senso sulle lettere che si hanno di fronte. Un ulteriore fattore di vulnerabilità è che la chiave è un oggetto fisico specifico, la cui trasmissione è più difficile da occultare rispetto, ad esempio, ad una parola-chiave.

1.2. Argomenti di matematica connessi

Questo metodo di cifratura può essere messo facilmente in relazione con le trasformazioni geometriche di rotazione e con le simmetrie e asimmetrie: la griglia nel suo complesso è asimmetrica per la presenza dei fori, ed è grazie a questa asimmetria che funziona il meccanismo di cifratura, tuttavia ha un ruolo fondamentale anche la simmetria del quadrato per rotazioni di un angolo retto, nonché il fatto che la scelta dei fori da praticare deve essere effettuata in modo da non forare celle che sono collegate dalle suddette rotazioni.

2. Cifratura per sostituzione con chiave

La cifratura per sostituzione, come suggerisce il nome stesso, consiste nel sostituire ogni lettera dell'alfabeto con un'altra in tutte le sue occorrenze. Per fare questo si affianca all'alfabeto normale (*in chiaro*) un secondo alfabeto, detto *alfabeto cifrante*, che non è altro che l'alfabeto normale in cui tutte (o quasi tutte) le lettere sono state cambiate di posizione. L'accostamento dell'alfabeto in chiaro con quello cifrante è detto *coppia cifrante*. Un esempio è

A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z
F	L	R	X	C	I	O	U	B	H	N	T	Z	E	K	Q	W	A	G	M	S	Y	D	J	P	V

Tabella 1: Coppia cifrante.

che trasforma la frase "Frequento il Liceo Scientifico A. SCACCHI" da

FREQUENTOILLICEOSCIENTIFICOASCACCHI

a

IACWSCEMKBTTBRCKGRBCEMBIBRKFGFRFRUB

La trasformazione è eseguita cercando le lettere del messaggio in chiaro nella prima riga (in verde) della coppia cifrante e sostituendole con la corrispondente lettera della seconda riga (in giallo).

Serve un meccanismo che permetta di generare alfabeti cifranti efficaci (cioè con poche o nulle corrispondenze tra lettere identiche) e che sia facile da comunicare. Questo può essere fatto mediante una singola parola (che dovrà essere comunicata al destinatario in modo indipendente dal messaggio cifrato): la *parola chiave*. Si procede nel modo seguente.⁴

1. Si sceglie una parola chiave, ad esempio NICOLA.
2. La si scrive in un foglio quadrettato e poi sotto ogni lettera si inserisce un numero che indica che posizione alfabetica ha quella lettera in relazione a tutte le altre lettere della parola chiave, ad esempio

⁴La procedura qui presentata per generale un alfabeto cifrante non è altro che un meccanismo rapido ed efficace di permutazione di una lista di simboli, in questo caso le lettere dell'alfabeto. Lo stesso meccanismo può essere usato direttamente sul messaggio, per mescolarne le lettere senza sostituirle, ottenendo quindi un altro metodo di crittografia per trasposizione oltre quello con griglia visto precedentemente. Naturalmente questo secondo metodo per trasposizione possiede esattamente gli stessi difetti di quello con griglia, dato che essi derivano dall'operazione generale di trasposizione. Quale possa essere più comodo da usare dipende dalla facilità con cui si può trasmettere in sicurezza la chiave al destinatario designato, se in forma di griglia o di parola chiave.

3.

N	I	C	O	L	A
5	3	2	6	4	1

sotto la A va il numero 1 perché la A è la prima lettera in ordine alfabetico tra quelle nella parola NICOLA, sotto la C va il numero 2 perché la C è la seconda lettera in ordine alfabetico tra quelle nella parola NICOLA, e così via fino ad inserire sotto la O il numero 6 perché la O è la sesta e ultima lettera in ordine alfabetico tra quelle nella parola NICOLA. Se la parola chiave contiene una lettera che si ripete, allora i numeri verranno messi in progressione da sinistra a destra sotto le varie lettere uguali, ad esempio se la parola chiave è CASSA i numeri verranno assegnati così

C	A	S	S	A
3	1	4	5	2

4. Successivamente scriviamo sotto alla riga dei numeri l'alfabeto secondo il normale ordine, procedendo da sinistra a destra e dall'alto in basso, inserendo ogni lettera sotto un numero. Per la parola chiave NICOLA si ha

N	I	C	O	L	A
5	3	2	6	4	1
A	B	C	D	E	F
G	H	I	J	K	L
M	N	O	P	Q	R
S	T	U	V	W	X
Y	Z				

Tabella 2: trasposizione dell'alfabeto con parola chiave.

mentre per la parola chiave CASSA si ha

C	A	S	S	A
3	1	4	5	2
A	B	C	D	E
F	G	H	I	J
K	L	M	N	O
P	Q	R	S	T
U	V	W	X	Y
Z				

5. Si copia la sequenza di lettere dall'alto in basso una colonna dopo l'altra, seguendo l'ordine delle colonne dato dai numeri in alto. Ad esempio dalla parola chiave NICOLA si ottiene la coppia cifrante vista all'inizio

A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z
F	L	R	X	C	I	O	U	B	H	N	T	Z	E	K	Q	W	A	G	M	S	Y	D	J	P	V

mentre dalla parola chiave CASSA si ottiene la coppia cifrante seguente

A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z
B	G	L	Q	V	E	J	O	T	Y	A	F	K	P	U	Z	C	H	M	R	W	D	I	N	S	X

Chi riceve il messaggio deve necessariamente conoscere la parola chiave per ricostruire la coppia cifrante, con cui può poi decifrare il messaggio: la decifrazione è eseguita cercando le lettere del messaggio cifrato nella seconda riga (in giallo) della coppia cifrante e sostituendole con la corrispondente lettera della prima riga (in verde).

2.1. Vantaggi e svantaggi di questo metodo

Rispetto alla crittografia per trasposizione qui non abbiamo più il problema della presenza delle stesse lettere originali, per cui questo metodo è inattaccabile da metodi combinatori sul messaggio.

Resta però la possibilità di indovinare alcune corrispondenze su base statistica: in ogni lingua ci sono lettere che sono più utilizzate delle altre,⁵ e quindi si possono provare ad individuare alcune corrispondenze vedendo quali sono i caratteri che compaiono più spesso, quelli un po' più rari, poi quelli ancora meno rari, e così via, in tutti i messaggi che si sanno essere stati cifrati con la stessa coppia cifrante.

⁵Ad esempio, in italiano le lettere E, A, I, O compaiono all'incirca lo stesso numero di volte e nettamente più di tutte le altre, segue poi il gruppo di lettere N, R, T, L, S, C D, poi quelle delle lettere U, P, M, V, G, poi quello di B, F, H, Q, Z e infine, molto raramente, il gruppo delle lettere J, K, W, X, Y.

Per ovviare a questo inconveniente si può passare alla crittografia per sostituzione polialfabetica che consiste nell'utilizzare insieme più di una parola chiave, ciascuna delle quali genera un alfabeto cifrante, e i diversi alfabeti cifranti vengono utilizzati alternativamente secondo un ordine prestabilito. Ad esempio supponiamo di usare come parole chiave NICOLA, CASSA e TURING. Gli alfabeti cifranti relativi alle prime due parole chiave li conosciamo già mentre quello generato da TURING è

T	U	R	I	N	G
5	6	4	2	3	1
A	B	C	D	E	F
G	H	I	J	K	L
M	N	O	P	Q	R
S	T	U	V	W	X
Y	Z				

Ora scriviamo i tre alfabeti cifranti uno sotto l'altro e sotto l'alfabeto in chiaro

A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z
F	L	R	X	C	I	O	U	B	H	N	T	Z	E	K	Q	W	A	G	M	S	Y	D	J	P	V
B	G	L	Q	V	E	J	O	T	Y	A	F	K	P	U	Z	C	H	M	R	W	D	I	N	S	X
F	L	R	X	D	J	P	V	E	K	Q	W	C	I	O	U	A	F	M	S	Y	B	H	N	T	Z

Tabella 3: Quaterna cifrante per trasposizione polialfabetica.

La cifratura ora procede in modo simile a prima con la differenza che non si usa lo stesso alfabeto cifrante per tutte le lettere, ma si scende di un rigo ad ogni lettera per poi tornare al primo alfabeto cifrante una volta arrivati all'ultimo rigo.

Nel nostro esempio abbiamo tre diversi alfabeti cifranti, allora le lettere nelle posizioni 1, 4, 7, 10, etc (in generale quelle di posto $3n \sim 2$) sono sostituite da lettere del primo alfabeto cifrante (rigo giallo), le lettere nelle posizioni 2, 5, 8, 11, etc (quelle di posto $3n \sim 1$) sono sostituite da lettere del secondo alfabeto cifrante (rigo rosso), le lettere nelle posizioni 3, 6, 9, 12, etc (quelle di posto $3n$) sono sostituite da lettere del terzo alfabeto cifrante. Pertanto il messaggio

FREQUENTOIL LICEOSCIENTIFICOASCACCHI

viene cifrato in

IHDWWDEROBFWB LDKMRBVIMTJBLOFMRFLRUT

Con questo metodo una stessa lettera del messaggio in chiaro viene cifrata con lettere diverse e una stessa lettera cifrante corrisponde a più lettere in chiaro, a seconda della posizione. Pertanto i ragionamenti su base statistica diventano possibili solo sulla scorta di un numero di messaggi, cifrati con lo stesso gruppo di chiavi, molto più elevato che nel caso di una coppia cifrante.

Tuttavia, come per la cifratura per trasposizione, se si ha a disposizione un computer abbastanza potente, la decifratura di messaggi cifrati in questo modo diventa abbastanza facile, perché è possibile provare tutte le possibili combinazioni di alfabeti cifranti in relativamente poco tempo.

Enigma, la macchina cifrante con cui i nazisti codificavano e decodificavano le proprie comunicazioni militari durante la Seconda Guerra Mondiale, utilizzava un sistema polialfabetico di questo tipo, ma era talmente complesso che nessun gruppo di esseri umani avrebbe potuto scoprire la chiave nelle 24 ore dopo le quali la chiave veniva nuovamente cambiata. Tuttavia la costruzione di una macchina calcolatrice automatica (un computer) da parte del matematico Alan Turing permise di risolvere questo problema in pochi minuti, fornendo agli Alleati uno strumento preziosissimo che in pochi anni portò alla fine della guerra con la sconfitta della Germania nazista.

2.2. Argomenti di matematica connessi

La prima e immediata connessione che si può analizzare è quella con il concetto di permutazione, cioè di riordinamento di un insieme: l'alfabeto in chiaro è una specifica permutazione, o ordinamento (scelto per convenzione come quello "giusto"), dell'insieme delle lettere dell'alfabeto; ogni alfabeto cifrante non è altro che una delle altre possibili permutazioni di tale insieme. Pertanto ci si possono porre i seguenti problemi:

1. *Quanti sono i possibili alfabeti cifranti?*

La risposta ovviamente è

$$26! - 1 = 403'291'461'126'605'635'583'999'999$$

cioè il numero di possibili permutazioni di un insieme di 26 elementi a cui sottraiamo 1 perché l'alfabeto in chiaro non va considerato.

2. *Quanti sono i possibili alfabeti cifranti realmente efficaci?*

Questa domanda ha senso perché se pensiamo ad un alfabeto cifrante semplicemente come ad un ordinamento diverso delle lettere dell'alfabeto, allora basta scambiare di posto anche solo due lettere per averne uno, ma tale alfabeto cifrante evidentemente non avrebbe alcuna efficacia, dato che il messaggio cifrato sarebbe sostanzialmente identico a quello in chiaro.

Ci si può porre quindi la seguente domanda: quante sono le permutazioni dell'alfabeto in cui nessuna lettera conserva il proprio posto rispetto all'ordinamento convenzionale? Queste permutazioni sono chiamate anche *dismutazioni* e, più precisamente, sono le permutazioni senza punti fissi. Dal calcolo combinatorio sappiamo che il numero di dismutazioni di n elementi è $|D_n| = \sum_{k=0}^n (-1)^k \frac{n!}{k!}$.⁶ Nel nostro caso dobbiamo quindi calcolare

$$|D_{26}| = \sum_{k=0}^{26} (-1)^k \frac{26!}{k!} = 148'362'637'348'470'135'821'287'825$$

che è poco più di un terzo del numero calcolato precedentemente. Il che vuol dire che scegliendo a caso le parole chiave, circa due volte su tre ci si imbatte in un alfabeto cifrante che lascia invariata almeno una lettera.

La crittografia per sostituzione richiama direttamente i concetti di funzione, funzione inversa e corrispondenza biunivoca: ogni coppia cifrante rappresenta una funzione che ha per dominio e codominio l'alfabeto e in particolare è una corrispondenza biunivoca, perché l'accoppiamento tra lettera in chiaro e lettera cifrata è unico. La corrispondenza biunivoca è condizione necessaria per poter poi procedere alla decifratura del messaggio: solo con una corrispondenza biunivoca è possibile invertire completamente l'accoppiamento tra le lettere e quindi decifrare il messaggio con chiarezza. Come detto prima, una coppia cifrante non è altro che una permutazione nell'insieme delle lettere dell'alfabeto, quindi questa funzione può essere descritta come $y = p(x)$ dove $x, y \in A = \{l | l \text{ è una lettera dell'alfabeto}\}$ e p è una delle permutazioni (o una delle dismutazioni) di A .

A questo punto si potrebbe obiettare che nella variante polialfabetica non c'è una corrispondenza biunivoca (alla stessa lettera in chiaro possono corrispondere più lettere cifrate e viceversa) eppure si riesce comunque a decifrare il messaggio. Questa obiezione, solo apparentemente giusta, permette di collegare questo metodo con il prodotto cartesiano di due insiemi.

La cifratura polialfabetica cifra le lettere in modo diverso in base alla posizione. Ciò significa che tale metodo equivale ad una funzione con dominio e codominio non nell'insieme A (come nel caso di singola coppia cifrante) ma nell'insieme $M = \{(x, k) | x \in A \wedge k \in J_L\} = A \times J_L$, con $J_L = \{1, \dots, L\}$ dove L è il numero di caratteri del messaggio (k indica la posizione della lettera x nel messaggio), perché per stabilire come si trasforma la lettera x si deve tenere in considerazione la sua posizione k nel messaggio. Stabiliti dominio e codominio, la corrispondenza è costruita nel modo seguente: si scelgono n permutazioni dell'insieme A (numerate da 0 a $n-1$), e poi si applica una permutazione al primo elemento della coppia (x, k) scegliendo la permutazione in base al valore di k :⁷

$$(y, k) = f(x, k) = (p_r(x), k) \wedge r = k \bmod n$$

⁶Il calcolo esplicito di questa formula è riportato in appendice.

⁷La scrittura $k \bmod n$ è solo un modo compatto per indicare il resto intero della divisione di $k : n$.

Per cui è vero che esistono lettere in chiaro (le x) uguali a cui vengono associate lettere cifrate (le y) diverse e viceversa, ma questo non contraddice la corrispondenza biunivoca perché guardando solo le lettere si sta guardando solo un "pezzo" degli enti messi in corrispondenza: gli elementi della corrispondenza sono le coppie lettera-posizione e non le sole lettere.

3. Cifratura mediante matrici

Il metodo presentato di seguito è una semplificazione dei metodi utilizzati realmente. La chiave di questo metodo è una sequenza di numeri interi di lunghezza pari ad un quadrato perfetto n^2 , in modo che la matrice $n \times n$ che si ottiene organizzando questi numeri al suo interno sia una matrice $\underline{\underline{M}}$ invertibile. Tale sequenza può anche essere rappresentata da un singolo numero che poi viene suddiviso in n^2 gruppi di cifre della stessa lunghezza.⁸ Ad esempio il numero 70500030001010100, suddiviso in blocchi

di due cifre a partire dalle unità, corrisponde alla matrice $\underline{\underline{M}} = \begin{pmatrix} 7 & 5 & 0 \\ 3 & 0 & 1 \\ 1 & 1 & 0 \end{pmatrix}$, la cui

inversa è $\underline{\underline{M}}^{-1} = \frac{1}{2} \begin{pmatrix} 1 & 0 & -5 \\ -1 & 0 & 7 \\ -3 & 2 & 15 \end{pmatrix}$.⁹

Successivamente si genera una sequenza di numeri interi convertendo il messaggio in una sequenza numerica mediante una corrispondenza predeterminata. In ambito informatico le più usate sono le tabelle ASCII e Unicode. Per semplicità, useremo il numero ordinale di ogni lettera, senza considerare punteggiatura, accenti, spazi e la distinzione tra maiuscole e minuscole. La lunghezza della sequenza deve essere un multiplo di n , se questa condizione non è soddisfatta allora si allunga la sequenza con un'opportuna quantità di 0. Ad esempio la frase "Frequento il Liceo Scientifico A. SCACCHI" viene trasformata nella sequenza

6, 18, 5, 17, 21, 5, 14, 20, 15, 9, 12, 12, 9, 3, 5, 15, 19, 3,

9, 5, 14, 20, 9, 6, 9, 3, 15, 1, 19, 3, 1, 3, 3, 8, 9, 0

Poi la sequenza viene suddivisa in vettori di lunghezza n

$(6, 18, 5), (17, 21, 5), (14, 20, 15), (9, 12, 12), (9, 3, 5), (15, 19, 3),$

$(9, 5, 14), (20, 9, 6), (9, 3, 15), (1, 19, 3), (1, 3, 3), (8, 9, 0)$

⁸Nel conteggio delle cifre possono eventualmente essere considerati anche degli 0 che occupano i posti di valore più alto.

⁹Una sintesi su matrici, determinanti e matrici inverse è riportata in appendice.

Infine ogni vettore viene trasformato moltiplicandolo per la matrice-chiave $\underline{\underline{M}}$

$(132, 23, 24), (224, 56, 38), (198, 57, 34), (123, 39, 21), (78, 32, 12), (200, 48, 34),$

$(88, 41, 14), (185, 66, 29), (78, 42, 12), (102, 6, 20), (22, 6, 4), (101, 24, 17)$

Il ricevente deve naturalmente conoscere la matrice $\underline{\underline{M}}$ per poter procedere alla decifrazione moltiplicando ogni vettore ricevuto per $\underline{\underline{M}}^{-1}$.

3.1 Vantaggi e svantaggi di questo metodo

Ciascun metodo di questa classe è molto facile da implementare direttamente in un programma, perché il calcolo matriciale coinvolge solo le operazioni elementari di somma e prodotto che sono previste in qualunque linguaggio di programmazione.¹⁰ Invece le operazioni relative ai classici metodi a chiave privata (quelli inventati prima dell'avvento dei computer) sono molto complessi da realizzare tramite software, in quanto la manipolazione diretta dei caratteri di stringhe testuali richiede operazioni superiori.

Questi metodi comunque soffrono della stessa vulnerabilità generale di tutti i sistemi a chiave privata: se usati da soli o insieme ad altri metodi a chiave privata, richiedono di trasmettere la chiave in chiaro al destinatario designato. Tuttavia essi possono essere accoppiati ad un sistema a chiave pubblica,¹¹ che viene usato per trasmettere la chiave privata. Di solito si procede in questo modo perché la crittografia a chiave pubblica richiede una grande quantità di calcoli che rallenterebbe di molto le comunicazioni se usata direttamente su messaggi lunghi.

3.2 Argomenti di matematica connessi

La connessione è ovvia: vettori e matrici, in particolare il significato, l'esistenza e il calcolo della matrice inversa di una matrice quadrata data.

Bibliografia

- [1] SACCO L. (autore), BONA VOGLIA P. (curatore), **Manuale di Crittografia**, *You-canprint*, 2014.

¹⁰Inoltre la maggioranza dei linguaggi di programmazione contengono routine già pronte per il calcolo con vettori e matrici.

¹¹come il cifrario RSA, che sarà esposto in un articolo successivo.

L'equazione differenziale di Karl Pearson

Karl Pearson's differential equation

Antonino Onorato¹

Abstract

This article presents Pearson curves' theory

1. Karl Pearson e la scuola biometrica

Nell'ambito degli studi di statistica metodologica, il nome di Karl Pearson è legato soprattutto alle nozioni di correlazione e di test statistico, con l'introduzione del coefficiente di correlazione r (1) e del test chi quadrato (2).

Questi fondamentali contributi sono sicuramente quelli più importanti e noti, ma non esauriscono il notevole apporto teorico e applicativo dato dal grande studioso inglese alla statistica metodologica.

Pearson aveva una personalità eclettica dai molteplici interessi culturali e scientifici. Si laurea in matematica nel 1879 all'Università di Cambridge, dove ebbe come docenti personalità del calibro di James Clerk Maxwell, George Gabriel Stokes, Arthur Cayley e Isaac Todhunter.

Successivamente, si trasferì in Germania per proseguire gli studi. Nelle Università di Heidelberg e Berlino studiò letteratura tedesca medievale, diritto, Metafisica, storia e cultura tedesca, storia delle religioni e filosofia. Lesse e scrisse su molti argomenti: poesia, teatro, etica, socialismo e diritti delle donne. Per un breve periodo esercitò anche la professione forense. Nel 1894, a ventisei anni, divenne professore di matematica applicata e meccanica all'Università college di Londra.

Inizialmente scettico e diffidente nell'applicazione dei metodi delle scienze esatte ai problemi di ereditarietà e di politica economica (3), Pearson cambiò opinione dopo l'incontro con Walter Frank Raphael Weldon.

¹ Socio dell'Unione Matematica Italiana. Responsabile di Sezione Consiglio di Stato-Tar.

Professore di zoologia e suo collega all'Università college di Londra, Weldon, con esempi concreti, tratti dall'analisi dei dati zoometrici riguardanti alcune misure realizzate sui crostacei decapodi (4), riuscì a dimostrare che le tecniche statistiche esposte da Galton nel suo "*Natural Inheritance* (5)"- Eredità Naturale- si potessero applicare con successo non solo in ambito antropologico, ma anche in quello biologico. Alle innegabili, notevoli capacità analitiche e teoriche, Pearson riuscì ad abbinare anche quelle di organizzatore e stimolatore. Nel 1901 con Galton e Weldon fondò la rivista *Biometrika* e l'anno dopo avviò presso l'University college di Londra un laboratorio di biometria con la finalità di dimostrare, attraverso l'applicazione dei nuovi strumenti della teoria statistica, le ipotesi della teoria evoluzionista. Nel 1920 unì il laboratorio di Eugenetica di Galton a quello di biometria creando, primo nel suo genere, il Dipartimento di Statistica Applicata (6).

2. Pearson quale creatore dei principali concetti della statistica descrittiva

Come accennato all'inizio del paragrafo precedente, Pearson introdusse il coefficiente di correlazione r e il test chi quadrato, ma questi importanti contributi, sicuramente i più noti e applicati, rappresentano però soltanto una parte della notevole produzione scientifica di concetti e metodi che egli sviluppò nell'ambito della statistica descrittiva.

Come ha acutamente rilevato Madrid Casado, "Pearson presentò molte delle sue idee ambiziose in una serie di conferenze serali che tenne tra il 1891 e il 1894 presso il Gresham College.

Le prime otto conferenze riguardarono gli aspetti fondamentali della filosofia della scienza, che furono raccolti nel libro *The Grammar of Science* (7), mentre le 30 conferenze restanti furono dedicate integralmente alla geometria della statistica e alla geometria del caso (8)". Il testo *The Grammar of Science* riassume l'epistemologia di Pearson nella quale si mescolano l'idealismo del filosofo Neokantiano Kuno Fisher e il positivismo di Ernst Mach (9).

La scienza, sostiene Pearson, deve limitarsi a descrivere i fatti osservabili, evitando qualsiasi tipo di ricaduta nella metafisica. Le leggi scientifiche non sono spiegazioni causali, ma riassunti ordinati dei fenomeni (10). Inoltre, fu un convinto sostenitore dell'unità della scienza, concepita però come unità di metodo e non di oggetto. Mentre le prime otto conferenze tenute da Pearson riguardarono gli aspetti fondamentali della sua filosofia della scienza, nelle successive trenta conferenze espose quelli che oggi si possono considerare i concetti fondamentali della statistica descrittiva. In particolare, in una di queste conferenze introdusse gli istogrammi, e poi sviluppò il concetto di deviazione standard che "a partire dal 1893 sostituì quello di errore probabile introdotto dall'astronomo Friedrich W. Bessel intorno al 1815 (11)". Se Quetelet e Galton rivalutarono, rispettivamente, la media e la mediana, Pearson tenne a battesimo la radice quadrata della media dei quadrati della differenza di ogni dato rispetto alla media, ossia la cosiddetta Deviazione Standard, indicata con il simbolo σ . Nelle conferenze successive Pearson, nell'ambito della sua attività di creatore e divulgatore di concetti statistici ideò, accanto alla deviazione standard, un altro indice di variabilità,

il coefficiente di variazione, e due misure descrittive di fondamentale importanza e precisamente il coefficiente di asimmetria e la curtosi.

3. Origine e sviluppo dell'equazione differenziale di Pearson

3.1 I presupposti che portarono a definire l'equazione

Riguardo al rapporto tra la curva normale e le curve anomale, non normali, Angiolo Maros dell'Oro, così si esprime: "Pearson, come Edgeworth, dapprima pensava che le curve anomale fossero sempre riconducibili a due o più curve normali."

Nel 1892 Weldon aveva trovato nei granchi di mare di Napoli, a proposito della loro ampiezza frontale, una curva a doppia "gobba" (bimodale) ed era riuscito a risolverla in due curve normali con transvariazione tra loro. Nell'autunno del 1893 Pearson aveva fatto qualcosa del genere, come risulta da una memoria presentata alla Royal Society e pubblicata pochi mesi più tardi come le prime delle sue "Contributions to the mathematical theory of evolution", dal titolo "on the dissection of asymmetrical frequency curves" (12). Ma a prescindere da questi risultati, il matematico inglese voleva comunque trovare un modo per interpretare i dati senza forzarne la normalizzazione, senza distorcere le forme delle curve di frequenza. Non voleva scartare la possibilità che esistesse una asimmetria reale nei dati di partenza. Queste considerazioni lo portarono a definire una equazione differenziale dalla quale, integrata con opportuni valori dei parametri, discende un'ampia famiglia di curve asimmetriche (13).

3.2 Sviluppi analitici

3.2.1 Fondamenti probabilistici dell'equazione differenziale

L'equazione differenziale del Pearson è definita dalla seguente espressione formale:

$$\frac{dy}{dx} = \frac{y(x+a)}{b_0 + b_1x + b_2x^2} \quad (3.1)$$

dove a è la distanza tra l'origine e la media, e b_0 , b_1 e b_2 sono parametri. Più precisamente, al denominatore viene indicato lo sviluppo in serie di Taylor, limitato al terzo termine, di una funzione $f(x)$. La struttura dell'equazione (3.1) esprime molto bene la dinamica di molte curve di frequenza e di probabilità che si incontrano spesso nelle analisi dei fenomeni reali. Tali curve hanno tutte alcune caratteristiche comuni: a partire da un valore iniziale della variabile indipendente x le loro ordinate y crescono via via dal valore zero fino ad un valore massimo, poi decrescono fino a ritornare al valore zero in corrispondenza ad un valore finale della variabile x e in tali punti le curve risultano essere tangenti all'asse delle ascisse. Ebbene, l'equazione (3.1) non fa

altro che esprimere in forma differenziale le curve che possiedono tali caratteristiche. Dunque, se Pearson avesse voluto, avrebbe potuto evitare manipolazioni dell'equazione (3.1) in quanto già di per sé fornisce una dimostrazione formale dell'attitudine a descrivere diversi tipi di distribuzioni di frequenza o probabilità, ma egli volle dare all'equazione (3.1) un fondamento probabilistico. Il modello teorico di riferimento è la distribuzione ipergeometrica o dell'estrazione in blocco. Ora, se da un'urna contenete N palline, di cui Np sono, ad esempio, bianche e Nq nere, con $p+q=1$, la probabilità $p(x)$ che estraendo un blocco di n palline x siano bianche è data proprio dalla distribuzione ipergeometrica.

$$f(x) = \frac{\left(\frac{Np}{x}\right)\left(\frac{Nq}{n-x}\right)}{\left(\frac{N}{n}\right)} \quad (3.2)$$

Vediamo attraverso quale procedimento sia possibile giungere a definire tale distribuzione. Il numero dei possibili blocchi di n palline estratte è $\left(\frac{N}{n}\right)$. Il numero dei casi favorevoli si ottiene considerando che le x palline bianche possono essere estratte in $\left(\frac{Np}{x}\right)$ modi diversi dalle Np palline bianche dell'urna, e analogamente le $n-x$ palline nere che completano il blocco in $\left(\frac{Nq}{n-x}\right)$ modi; inoltre, per ottenere un blocco si possono associare in tutti i modi possibili le x palline bianche con le $n-x$ palline nere. Quindi, il numero di blocchi diversi contententi x palline bianche è $\left(\frac{Np}{x}\right)\left(\frac{Nq}{n-x}\right)$, e la probabilità è data dalla (3.2).

Tenendo presente la (3.1) ed essendo la distribuzione ipergeometrica discreta, possiamo scrivere:

$$\frac{\Delta y}{y \Delta x} = \frac{y_x - y_{x-1}}{\frac{1}{2}(y_x + y_{x-1})} \quad (3.3)$$

Nella quale è:

$$y(x) = \frac{\left(\frac{Np}{x}\right)\left(\frac{Nq}{n-x}\right)}{\left(\frac{N}{n}\right)} \quad (3.4)$$

e

$$y(x-1) = \frac{\left(\frac{Np}{x-1}\right)\left(\frac{Nq}{n-x+1}\right)}{\left(\frac{N}{n}\right)} \quad (3.5)$$

Sottraendo la (3.5) dalla (3.4) si ha:

$$y_x - y_{x-1} = \frac{1}{\frac{N}{n}} \left\{ \frac{Np(Np-1)\dots(Np-x+1)}{x!} \frac{Nq(Nq-1)\dots(Nq-n+x+1)}{(n-x)!} - \frac{Np(Np-1)\dots(Np-x+2)}{(x-1)!} \frac{Nq(Nq-1)\dots(Nq-n+x)}{(n-x+1)!} \right\} \quad (3.6)$$

Tenendo presente che $x! = (x-1)!x$ e $(n-x+1)! = (n-x)!(n-x+1)$, possiamo quindi mettere in evidenza le espressioni $\frac{(Np-x+1)}{x}$ e $\frac{(Nq-n+x)}{(n-x+1)}$, rispettivamente nel primo e nel secondo termine della (3.6), e scrivere l'espressione nel seguente modo:

$$y_x - y_{x-1} = \frac{1}{\frac{N}{n}} \frac{Np(Np-1)\dots(Np-x+2)}{(x-1)!} \frac{Nq(Nq-1)\dots(Nq-n+x+1)}{(n-x)!} \left[\frac{Np-x+1}{x} - \frac{Nq-n+x}{(n-x+1)} \right] \quad (3.7)$$

Sostituendo nella (3.3) si ha:

$$\frac{y_x - y_{x-1}}{\frac{1}{2}(y_x + y_{x-1})} = 2 \frac{\frac{Np-x+1}{x} - \frac{Nq-n+x}{n-x+1}}{\frac{Np-x+1}{x} + \frac{Nq-n+x}{n-x+1}} \quad (3.8)$$

$$= 2 \frac{(Np-x+1)(n-x+1) - x(Nq-n+x)}{(Np-x+1)(n-x+1) + x(Nq-n+x)} \quad (3.9)$$

$$= 2 \frac{(Nnp + Np + n + 1) - x(N + 2)}{(Nnp + Np + n + 1) - x[N(p-q) + 2(n+1)]} \quad (3.10)$$

$$= 2 \frac{(Np+1)(n+1) - x(N+2)}{(Np+1)(n+1) - x[N(p-q) + 2(n+1)] + 2x^2} \quad (3.11)$$

Ora, dividendo numeratore e denominatore della (3.11) per $-2(N+2)$ e semplificando si ha:

$$= \frac{-\frac{(Np+1)(n+1)}{(N+2)} + x}{-\frac{(Np+1)(n+1)}{2(N+2)} + \frac{[N(p-q)+2(n+1)]}{2(N+2)}x - \frac{1}{(N+2)}x^2} \quad (3.12)$$

Ponendo poi:

$$\begin{aligned} -\frac{(Np+1)(n+1)}{(N+2)} &= a \\ -\frac{(Np+1)(n+1)}{2(N+2)} &= b_0 \\ \frac{N(p-q)+2(n+1)}{2(N+2)} &= b_1 \\ -\frac{1}{(N+2)} &= b_2 \end{aligned} \quad (3.13)$$

Si ha:

$$\frac{a+x}{b_0+b_1x+b_2x^2} \quad (3.14)$$

Che corrisponde al secondo membro dell'equazione differenziale:

$$\frac{dy}{ydx} = \frac{x+a}{b_0+b_1x+b_2x^2} \quad (3.15)$$

Che corrisponde a sua volta all'equazione (3.1). In tal modo, Pearson dimostra il fondamento probabilistico della sua equazione.

3.2.2 Parametri dell'equazione differenziale espressi in funzione dei momenti

La (3.1) può essere scritta anche in un altro modo, esprimendo i parametri in funzione dei momenti. Al riguardo si ha:

$$(b_0+b_1x+b_2x^2)\frac{dy}{dx} = y(x+a) \quad (3.16)$$

Moltiplicando primo e secondo membro per x^n ed integrando rispetto ad x si ha:

$$\int_{-\infty}^{+\infty} x^n (b_0 + b_1 x + b_2 x^2) \frac{dy}{dx} dx = \int_{-\infty}^{+\infty} x^n (x + a) y dx \quad (3.17)$$

Integrando per parti il primo membro tra i limiti indicati si ha:

$$\begin{aligned} \int_{-\infty}^{+\infty} [x^n (b_0 + b_1 x + b_2 x^2) y] - \int_{-\infty}^{+\infty} [nb_0 x^{n-1} + (n+1)b_1 x^n + (n+2)b_2 x^{n+1}] y dx \\ = \int_{-\infty}^{+\infty} x^{n+1} y dx + a \int_{-\infty}^{+\infty} x^n y dx \end{aligned} \quad (3.18)$$

Se l'espressione $[x^n (b_0 + b_1 x + b_2 x^2) y]$ si annulla agli estremi dell'intervallo di integrazione si ha:

$$\begin{aligned} - \int_{-\infty}^{+\infty} [nb_0 x^{n-1} + (n+1)b_1 x^n + (n+2)b_2 x^{n+1}] y dx = \\ = \int_{-\infty}^{+\infty} x^{n+1} y dx + a \int_{-\infty}^{+\infty} x^n y dx \end{aligned} \quad (3.19)$$

Che possiamo scrivere anche nel seguente modo:

$$\begin{aligned} a \int_{-\infty}^{+\infty} x^n y dx + \int_{-\infty}^{+\infty} nb_0 x^{n-1} y dx + (n+1) \int_{-\infty}^{+\infty} b_1 x^n y dx + \\ + \int_{-\infty}^{+\infty} (n+2) b_2 x^{n+1} y dx = - \int_{-\infty}^{+\infty} x^{n+1} y dx \end{aligned} \quad (3.20)$$

Tenendo presente che $\mu_n = \int_{-\infty}^{+\infty} x^n y dx$ indica il momento di ordine n dall'origine, la (3.20) assume la forma:

$$a\mu_n + nb_0\mu_{n-1} + (n+1)b_1\mu_n + (n+2)b_2\mu_{n+1} = -\mu_{n+1} \quad (3.21)$$

Ora, dando a n i valori 0,1,2,3 otteniamo un sistema di quattro equazioni che permette di determinare i valori dei quattro parametri che figurano nella (3.1). Il sistema che si sviluppa è:

$$\begin{cases} a\mu_0 + 0b_0 + b_1\mu_0 + 2b_2\mu_1 = -\mu_1 \\ a\mu_1 + b_0\mu_0 + 2b_1\mu_1 + 3b_2\mu_2 = -\mu_2 \\ a\mu_2 + 2b_0\mu_1 + 3b_1\mu_2 + 4b_2\mu_3 = -\mu_3 \\ a\mu_3 + 3b_0\mu_2 + 4b_1\mu_3 + 5b_2\mu_4 = -\mu_4 \end{cases} \quad (3.22)$$

Per semplificare il sistema, consideriamo i momenti dalla media aritmetica. Poiché $\bar{\mu}_0 = 1$ e $\bar{\mu}_1 = 0$, il sistema si riduce al seguente:

$$\begin{cases} a + b_1 = 0 \\ b_0 + 3b_2\bar{\mu}_2 = -\bar{\mu}_2 \\ a\bar{\mu}_2 + 3b_1\bar{\mu}_2 + 4b_2\bar{\mu}_3 = -\bar{\mu}_3 \\ a\bar{\mu}_3 + 3b_0\bar{\mu}_2 + 4b_1\bar{\mu}_3 + 5b_2\bar{\mu}_4 = -\bar{\mu}_4 \end{cases} \quad (3.23)$$

Risolvendo il sistema otteniamo i seguenti valori dei parametri:

$$\begin{aligned} a &= \frac{\bar{\mu}_3(\bar{\mu}_4 + 3\bar{\mu}_2^2)}{10\bar{\mu}_2\bar{\mu}_4 - 12\bar{\mu}_3^2 - 18\bar{\mu}_2^3} \\ b_0 &= \frac{\bar{\mu}_2(4\bar{\mu}_2\bar{\mu}_4 - 3\bar{\mu}_3^2)}{12\bar{\mu}_3^2 + 18\bar{\mu}_2^3 - 10\bar{\mu}_2\bar{\mu}_4} \\ b_1 &= \frac{-\bar{\mu}_3(\bar{\mu}_4 + 3\bar{\mu}_2^2)}{10\bar{\mu}_2\bar{\mu}_4 - 12\bar{\mu}_3^2 - 18\bar{\mu}_2^3} \\ b_2 &= \frac{2\bar{\mu}_2\bar{\mu}_4 - 3\bar{\mu}_3^2 - 6\bar{\mu}_2^3}{12\bar{\mu}_3^2 + 18\bar{\mu}_2^3 - 10\bar{\mu}_2\bar{\mu}_4} \end{aligned} \quad (3.24)$$

Sostituendoli nell'equazione (3.1) e moltiplicando numeratore e denominatore per -1, l'equazione (3.1) assume la nuova forma:

$$\frac{1}{y} \frac{dy}{dx} = - \frac{x + \frac{\bar{\mu}_3(\bar{\mu}_4 + 3\bar{\mu}_2^2)}{10\bar{\mu}_2\bar{\mu}_4 - 12\bar{\mu}_3^2 - 18\bar{\mu}_2^3}}{\frac{\bar{\mu}_2(4\bar{\mu}_2\bar{\mu}_4 - 3\bar{\mu}_3^2)}{10\bar{\mu}_2\bar{\mu}_4 - 12\bar{\mu}_3^2 - 18\bar{\mu}_2^3} + \frac{\bar{\mu}_3(\bar{\mu}_4 + 3\bar{\mu}_2^2)}{10\bar{\mu}_2\bar{\mu}_4 - 12\bar{\mu}_3^2 - 18\bar{\mu}_2^3} x + \frac{2\bar{\mu}_2\bar{\mu}_4 - 3\bar{\mu}_3^2 - 6\bar{\mu}_2^3}{10\bar{\mu}_2\bar{\mu}_4 - 12\bar{\mu}_3^2 - 18\bar{\mu}_2^3}} \quad (3.25)$$

Ponendo $\beta_1 = \frac{\bar{\mu}_3^2}{\bar{\mu}_2^3}$ e $\beta_2 = \frac{\bar{\mu}_4}{\bar{\mu}_2^2}$. L'equazione (3.25) si trasforma ulteriormente nel seguente modo:

$$\frac{1}{y} \frac{dy}{dx} = \frac{x + \frac{\sqrt{\mu_2} \sqrt{\beta_1(\beta_2+3)}}{2(5\beta_2-6\beta_1-9)}}{\frac{\bar{\mu}_2(4\beta_2-3\beta_1) + \sqrt{\mu_2} \sqrt{\beta_1(\beta_2+3)}x + (2\beta_2-3\beta_1-6)x^2}{2(5\beta_2-6\beta_1-9)}} \quad (3.26)$$

L'equazione (3.1), che permette di passare alla relazione $y = f(x)$, dipende dalle radici dell'equazione che sta al denominatore e queste, a loro volta, dipendono dal valore del discriminante ($b_1^2 - 4b_0b_2$), ossia dal valor del rapporto $\frac{b_1^2}{4b_0b_2}$ il quale, esprimendo i parametri in funzione di β_1 e β_2 , diventa:

$$F = \frac{\beta_1(\beta_2 + 3)^2}{4(4\beta_2 - 3\beta_1)(2\beta_2 - 3\beta_1 - 6)} \quad (3.27)$$

F è la cosiddetta funzione critica la quale, unitamente ai valori di β_1 e β_2 , consente di determinare la funzione teorica che meglio si adatta ai dati empirici.

3.2.3 Le principali funzioni di Pearson derivanti dall'integrazione della sua equazione differenziale

Le principali funzioni di Pearson, indicate dal nostro autore come curve del tipo I, del tipo II fino a quelle del tipo VII, più la curva normale, che non viene caratterizzata però come curva del tipo VIII, ma indicata con il proprio nome, si ottengono integrando l'equazione differenziale (3.1) tenendo conto del rapporto $\frac{b_1^2}{4b_0b_2}$ e della funzione critica (3.27). Vediamo nel dettaglio gli sviluppi analitici che caratterizzano le diverse funzioni.

TIPO I

Nelle funzioni del tipo I si ha $F < 0$. Pertanto, si hanno due radici reali e distanti. Se infatti la funzione critica:

$$F = \frac{\beta_1(\beta_2 + 3)^2}{4(4\beta_2 - 3\beta_1)(2\beta_2 - 3\beta_1 - 6)} = \frac{b_1^2}{4b_0b_2} < 0 \quad (3.28)$$

poiché il numeratore è sempre positivo, il segno verrà determinato dal denominatore, $4b_0b_2$, che è negativo. In tal caso l'espressione $\sqrt{b_1^2 - 4b_0b_2}$ è positiva e ci consente di definire le due radici reali e distinte. Dunque, si ha:

$$\frac{1}{y} \frac{dy}{dx} = \frac{x + a}{b_0 + b_1x + b_2x^2} \quad (3.29)$$

Indicando le radici del polinomio $b_0 + b_1x + b_2x^2$ nel seguente modo $-\alpha_1, \alpha_2$ ($\alpha_1 > 0, \alpha_2 > 0$), l'equazione assume allora la forma:

$$\frac{1}{y} \frac{dy}{dx} = \frac{1}{b_2} \frac{x+a}{(x+\alpha_1)(x-\alpha_2)} \quad (3.30)$$

Si ha:

$$\frac{1}{y} \frac{dy}{dx} = \frac{1}{b_2} \left[\frac{\alpha_1 - a}{\alpha_1 + \alpha_2} \frac{1}{x + \alpha_1} + \frac{\alpha_2 + a}{\alpha_1 + \alpha_2} \frac{1}{x - \alpha_2} \right]$$

Integrando si ha:

$$\log y = \frac{1}{b_2} \left[\frac{\alpha_1 - a}{\alpha_1 + \alpha_2} \log(x + \alpha_1) + \frac{\alpha_2 + a}{\alpha_1 + \alpha_2} \log(x - \alpha_2) \right] + c \quad (3.31)$$

Eliminando i logaritmi e ponendo $y_1 = e^c$ si ha:

$$y = y_1 (x + \alpha_1)^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 + \alpha_2}} (x - \alpha_2)^{\frac{1}{b_2} \frac{\alpha_2 + a}{\alpha_1 + \alpha_2}} \quad (3.32)$$

Con un cambiamento d'origine, che possiamo ottenere sostituendo ax l'espressione $x - a$, la funzione assume la forma:

$$y = y_1 (x - a + \alpha_1)^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 + \alpha_2}} (x - a - \alpha_2)^{\frac{1}{b_2} \frac{\alpha_2 + a}{\alpha_1 + \alpha_2}} \quad (3.33)$$

Che possiamo trasformare nel seguente modo:

$$y = y_1 \left[\left(1 + \frac{x}{\alpha_1 - a} \right) (\alpha_1 - a) \right]^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 + \alpha_2}} \left\{ \left(1 - \frac{x}{\alpha_2 + a} \right) [-(\alpha_2 + a)] \right\}^{\frac{1}{b_2} \frac{\alpha_2 + a}{\alpha_1 + \alpha_2}} \quad (3.34)$$

$$y = y_1 \left(1 + \frac{x}{\alpha_1 - a} \right)^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 + \alpha_2}} (\alpha_1 - a)^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 + \alpha_2}} \left(1 - \frac{x}{\alpha_2 + a} \right)^{\frac{1}{b_2} \frac{\alpha_2 + a}{\alpha_1 + \alpha_2}} (-\alpha_2 - a)^{\frac{1}{b_2} \frac{\alpha_2 + a}{\alpha_1 + \alpha_2}} \quad (3.35)$$

Ed infine, dopo aver posto

$$y_0 = y_1 (\alpha_1 - a)^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 + \alpha_2}} (-\alpha_2 - a)^{\frac{1}{b_2} \frac{\alpha_2 + a}{\alpha_1 + \alpha_2}}$$

$$\alpha_1 - a = a_1$$

$$\alpha_2 - a = a_2$$

$$v = \frac{1}{b_2} \frac{1}{\alpha_1 + \alpha_2}$$

Otteniamo la funzione tipo I nella forma indicata da Pearson, ossia:

$$y = y_0 \left(1 + \frac{x}{a_1}\right)^{va_1} \left(1 - \frac{x}{a_2}\right)^{va_2} \quad (3.36)$$

TIPO II

Le curve del secondo tipo si ottengono dalla (3.36) ponendo $a_1 = a_2 = a$. Si ha quindi

$$y = y_0 \left(1 + \frac{x}{a}\right)^{va} \left(1 - \frac{x}{a}\right)^{va}$$

$$y = y_0 \left[\left(1 + \frac{x}{a}\right) \left(1 - \frac{x}{a}\right)\right]^{va}$$

$$y = y_0 \left(1 - \frac{x^2}{a^2}\right)^m \quad (3.37)$$

Avendo posto $va = m$, la (3.37) rappresenta la famiglia delle funzioni del tipo II.

TIPO III

In questo caso nella (3.27) l'espressione al denominatore $2\beta_2 - 3\beta_1 - 6$ diventa nulla e di conseguenza la funzione critica assume un valore infinito. Nell'equazione differenziale (3.1) si ha quindi $b_2 = 0$ e la stessa assume la seguente forma:

$$\frac{dy}{dx} = \frac{y(x+a)}{b_0 + b_1 x} \quad (3.38)$$

Integrando si ha

$$\begin{aligned}
 \int \frac{1}{y} dy &= \int \frac{x+a}{b_0+b_1x} dx + c \\
 \log y &= \int \left(\frac{1}{b_1} + \frac{a-\frac{b_0}{b_1}}{b_1x+b_0} \right) dx + c \\
 \log y &= \int \frac{1}{b_1} dx + \int \left(\frac{a-\frac{b_0}{b_1}}{b_1x+b_0} \right) dx + c \\
 \log y &= \frac{x}{b_1} + \left(a - \frac{b_0}{b_1} \right) \frac{1}{b_1} \log(b_1x+b_0) + c \\
 y &= e^{\frac{x}{b_1} + \left(a - \frac{b_0}{b_1} \right) \frac{1}{b_1} \log(b_1x+b_0) + c} \\
 y &= y_1 e^{\frac{x}{b_1}} (b_1x+b_0)^{\left(a - \frac{b_0}{b_1} \right) \frac{1}{b_1}}
 \end{aligned} \tag{3.39}$$

Dove abbiamo posto $e^c = y_1$. Trasformiamo la forma della funzione ottenuta secondo la modalità indicata da Pearson. Per fare questo apportioniamo un cambiamento di origine della funzione sostituendo a x il valore $x + \left(a - \frac{b_0}{b_1} \right)$; si ha quindi:

$$\begin{aligned}
 y &= y_1 e^{\frac{1}{b_1} \left[x + \left(a - \frac{b_0}{b_1} \right) \right]} \left[b_1 \left(x + a - \frac{b_0}{b_1} \right) + b_0 \right]^{\left(a - \frac{b_0}{b_1} \right) \frac{1}{b_1}} \\
 y &= y_1 e^{\frac{1}{b_1} x + \frac{1}{b_1} a - \frac{b_0}{b_1^2}} [b_1x + b_1a]^{\left(a - \frac{b_0}{b_1} \right) \frac{1}{b_1}} \\
 y &= y_1 e^{\frac{1}{b_1} x + a \left(\frac{1}{b_1} - \frac{b_0}{ab_1^2} \right)} [b_1x + b_1a]^{\left(a - \frac{b_0}{b_1} \right) \frac{1}{b_1}} \\
 y &= y_1 e^{\frac{1}{b_1} x + a \left(\frac{1}{b_1} - \frac{b_0}{ab_1^2} \right)} \left[1 + \frac{x}{a} \right]^{a \left(\frac{1}{b_1} - \frac{b_0}{ab_1^2} \right)} [ab_1]^a \left(\frac{1}{b_1} - \frac{b_0}{ab_1^2} \right)
 \end{aligned}$$

Ponendo

$$y_0 = y_1 [ab_1]^a \left(\frac{1}{b_1} - \frac{b_0}{ab_1^2} \right)$$

e $\gamma = \frac{1}{b_1} - \frac{b_0}{ab_1^2}$, si ha:

$$y = y_0 e^{\frac{1}{b_1} x - a\gamma} \left(1 + \frac{x}{a} \right)^{a\gamma}$$

$$y = y_0 e^{\frac{1}{b_1} x} \left(1 + \frac{x}{a}\right)^{a\gamma}$$

Dove il nuovo valore y_0 è $y_0 e^{-a\gamma}$ poiché $\gamma = \frac{1}{b_1} - \frac{b_0}{ab_1^2}$, si ha:

$$\frac{1}{b_1} = \gamma + \frac{b_0}{ab_1^2}$$

Sostituendo nell'equazione si ha:

$$y = y_1 e^{x\left(\gamma + \frac{b_0}{ab_1^2}\right)} \left(1 + \frac{x}{a}\right)^{a\gamma}$$

Ponendo

$$\gamma + \frac{b_0}{ab_1^2} = -\gamma \Rightarrow \gamma + \frac{b_0}{ab_1^2} = -\left(\frac{1}{b_1} - \frac{b_0}{ab_1^2}\right)$$

Quindi

$$\gamma + \frac{b_0}{ab_1^2} = \left(-\frac{1}{b_1} + \frac{b_0}{ab_1^2}\right) \Rightarrow \gamma = -\frac{1}{b_1}$$

e

$$\frac{1}{b_1} = -\gamma$$

Con quest'ultima sostituzione, l'equazione (3.39) assume la forma indicata da Pearson:

$$y = y_0 \left(1 + \frac{x}{a}\right)^{\gamma a} e^{-\gamma x}$$

TIPO IV

In questo caso le radici dell'equazione $b_2 x^2 + b_1 x + b_0 = 0$ sono complesse. Integrando l'equazione (3.1) si ha:

$$\log y = \int \frac{x + a}{b_0 + b_1 x + b_2 x^2} dx \quad (3.40)$$

Scriviamo l'integrale al secondo membro nel seguente modo:

$$\int \frac{x + \frac{b_1}{2b_2} + a - \frac{b_1}{2b_2}}{b_2 \left[\left(x^2 + \frac{2xb_1}{2b_2} + \frac{b_1^2}{4b_2^2} \right) + \frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2} \right]} dx \quad (3.41)$$

Poniamo:

$$t = x + \frac{b_1}{2b_2}, a^2 = \frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2}, c = a - \frac{b_1}{2b_2} \quad (3.42)$$

Sostituendo nella (3.41) si ha:

$$\int \frac{t + c}{b_2 [t^2 + a^2]} dt$$

Che possiamo scrivere nel seguente modo:

$$\int \frac{t}{b_2 [t^2 + a^2]} dt + \int \frac{c}{b_2 [t^2 + a^2]} dt$$

Integrando si ha:

$$\frac{1}{2b_2} \log (t^2 + a^2) + \frac{c}{b_2 a} \arctan \frac{t}{a} + c$$

Che possiamo esprimere in funzione di x , con le sostituzioni (3.42), nel seguente modo:

$$\frac{1}{2b_2} \log \left[\left(x + \frac{b_1}{2b_2} \right)^2 + \frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2} \right] + \frac{a - \frac{b_1}{2b_2}}{b_2 \sqrt{\frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2}}} \arctan \frac{x + \frac{b_1}{2b_2}}{\sqrt{\frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2}}} + c$$

Tenendo conto che al primo membro si ha $\log y(x)$, la funzione soluzione è:

$$y(x) = y_1 \left[\left(x + \frac{b_1}{2b_2} \right)^2 + \frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2} \right]^{\frac{1}{2b_2}} e^{\frac{a - \frac{b_1}{2b_2}}{b_2 \sqrt{\frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2}}} \arctan \frac{x + \frac{b_1}{2b_2}}{\sqrt{\frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2}}}} \quad (3.43)$$

Sostituendo a x , $x - \frac{b_1}{2b_2}$, si ha:

$$y(x) = y_1 \left[x^2 + \frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2} \right]^{\frac{1}{2b_2}} e^{\frac{a - \frac{b_1}{2b_2}}{b_2 \sqrt{\frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2}}} \arctan \frac{x}{a}}$$

$$y(x) = y_1 \left[x^2 + \frac{b_0}{b_2} - \frac{b_1^2}{4b_2^2} \right]^{\frac{1}{2b_2}} e^{\frac{a - \frac{b_1}{2b_2}}{ab_2} \arctan \frac{x}{a}}$$

$$y(x) = y_1 [x^2 + a^2]^{\frac{1}{2b_2}} e^{\frac{a - \frac{b_1}{2b_2}}{ab_2} \arctan \frac{x}{a}}$$

Trasformando l'esponente $\frac{a - \frac{b_1}{2b_2}}{ab_2}$ nel seguente modo $-\left(\frac{b_1}{2b_2a} - \frac{1}{b_2}\right)$ e ponendo $\left(\frac{b_1}{2b_2a} - \frac{1}{b_2}\right) = v$ e $m = -\frac{1}{2b_2}$ si ha:

$$y = y_1 \frac{1}{(x^2 + a^2)^m} e^{-\text{varctang} \frac{x}{a}}$$

Ed ancora, scrivendo il termine $(x^2 + a^2)$ nel seguente modo $\left(1 + \frac{x^2}{a^2}\right) a^2$, si ha:

$$y = \frac{y_1}{\left[\left(1 + \frac{x^2}{a^2}\right) a^2\right]^m} e^{-\text{varctang} \frac{x}{a}}$$

Ed infine, ponendo $y_1 a^{-2m} = y_0$ otteniamo la funzione nella forma indicata da Pearson

$$y = \frac{y_0}{\left(1 + \frac{x^2}{a^2}\right)^m} e^{-\text{varctang} \frac{x}{a}} \quad (3.44)$$

TIPO V

In questo caso l'equazione (3.1), tenendo conto del fatto che le radici dell'equazione $b_2x^2 + b_1x + b_0 = 0$ sono reali ed uguali, avrà la forma:

$$\frac{dy}{dx} \frac{1}{y} = \frac{x + a}{b_2 \left(x + \frac{b_1}{2b_2}\right)^2} \quad (3.45)$$

Integrando si ha:

$$\begin{aligned} \int \frac{1}{y} dy &= \int \frac{1}{b_2} \frac{x + a}{\left(x + \frac{b_1}{2b_2}\right)^2} dx \\ \log y &= \int \frac{1}{b_2} \frac{dx}{\left(x + \frac{b_1}{2b_2}\right)} + \int \frac{a - \frac{b_1}{2b_2}}{b_2 \left(x + \frac{b_1}{2b_2}\right)^2} dx \end{aligned}$$

ossia

$$\begin{aligned} \log y &= \frac{1}{b_2} \log \left(x + \frac{b_1}{2b_2}\right) - \frac{a - \frac{b_1}{2b_2}}{b_2 \left(x + \frac{b_1}{2b_2}\right)} + c \\ y &= y_1 \left(x + \frac{b_1}{2b_2}\right)^{\frac{1}{b_2}} e^{-\frac{a - \frac{b_1}{2b_2}}{b_2 \left(x + \frac{b_1}{2b_2}\right)}} \end{aligned} \quad (3.46)$$

Dove abbiamo posto $e^c = y_1$ anche in questo caso con il cambiamento d'origine, ponendo al posto di x l'espressione $x - \frac{b_1}{2b_2}$, si ha:

$$y = y_1 x^{\frac{1}{b_2}} e^{-\frac{a - \frac{b_1}{2b_2}}{x b_2}}$$

Ponendo poi:

$$p = -\frac{1}{b_2} \text{ e } \gamma = \frac{a - \frac{b_1}{2b_2}}{b_2}$$

Si ha la funzione nella forma di Pearson:

$$y = y_0 x^{-p} e^{-\frac{\gamma}{x}} \quad (3.47)$$

TIPO VI

In questo caso le radici dell'equazione: dell'equazione $b_2 x^2 + b_1 x + b_0 = 0$ sono uguali e dello stesso segno. L'equazione (3.1) può essere scritta nel seguente modo:

$$\frac{1}{y} \frac{dy}{dx} = \frac{1}{b_2} \frac{x + a}{(x + \alpha_1)(x + \alpha_2)} \quad (3.48)$$

Ed ancora

$$\frac{1}{y} \frac{dy}{dx} = \frac{1}{b_2} \left(\frac{\alpha_1 - a}{\alpha_1 - \alpha_2} \frac{1}{x + \alpha_1} + \frac{a - \alpha_2}{\alpha_1 - \alpha_2} \frac{1}{x + \alpha_2} \right)$$

Integrando si ha

$$\begin{aligned} \log y &= \frac{1}{b_2} \int \left(\frac{\alpha_1 - a}{\alpha_1 - \alpha_2} \frac{1}{x + \alpha_1} + \frac{a - \alpha_2}{\alpha_1 - \alpha_2} \frac{1}{x + \alpha_2} \right) dx \\ \log y &= \frac{1}{b_2} \left[\frac{\alpha_1 - a}{\alpha_1 - \alpha_2} \log(x + \alpha_1) + \frac{a - \alpha_2}{\alpha_1 - \alpha_2} \log(x + \alpha_2) \right] + c \\ y &= y_1 (x + \alpha_1)^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 - \alpha_2}} (x + \alpha_2)^{\frac{1}{b_2} \frac{a - \alpha_2}{\alpha_1 - \alpha_2}} \end{aligned} \quad (3.49)$$

Avendo posto $e^c = y_1$. Cambiando l'origine della (3.49) si ha, ponendo al posto di x l'espressione $x - \alpha_2$,

$$y = y_1 (x - \alpha_2 + \alpha_1)^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 - \alpha_2}} x^{\frac{1}{b_2} \frac{a - \alpha_2}{\alpha_1 - \alpha_2}}$$

$$y = y_1 [x - (\alpha_2 - \alpha_1)]^{\frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 - \alpha_2}} x^{-\frac{1}{b_2} \frac{\alpha_2 - a}{\alpha_1 - \alpha_2}}$$

Ponendo $\alpha_2 - \alpha_1 = a_1$

e

$$\begin{cases} \frac{1}{b_2} \frac{\alpha_1 - a}{\alpha_1 - \alpha_2} = m_1 \\ \frac{1}{b_2} \frac{\alpha_2 - a}{\alpha_1 - \alpha_2} = m_2 \end{cases}$$

Si ha infine la funzione nella forma di Pearson:

$$y = y_0 (x - a_1)^{m_1} x^{-m_2} \quad (3.50)$$

TIPO VII

Questo tipo rappresenta un caso particolare del tipo IV, cioè si verifica quando $v=0$. La funzione è la seguente:

$$y = y_0 \left(1 + \frac{x^2}{a^2}\right)^{-m}$$

CURVA NORMALE

A differenza delle altre funzioni classificate da Pearson nell'ordine dal primo al settimo tipo, quella che rappresenta la curva normale viene inserita nel sistema delle curve di Pearson, ma non viene classificata nella tipologia "tipo", ma indicata con il proprio nome.

In questo caso nell'equazione (3.1) i coefficienti b_1 e b_2 del polinomio al denominatore del secondo membro sono nulli. L'equazione (3.1) assume quindi la forma:

$$\frac{1}{y} \frac{dy}{dx} = \frac{x + a}{b_0}$$

Integrando si ha:

$$\log y = \int \frac{a}{b_0} dx + \int \frac{x}{b_0} dx$$

Ossia, integrando e ponendo $c = \frac{a^2}{2b_0} + c_1$, si ha:

$$y = y_1 e^{\frac{x^2 + 2ax + a^2}{2b_0}}$$

Dove $y_1 = e^{c_1}$, si ha quindi:

$$y = y_1 e^{\frac{(x+a)^2}{2b_0}}$$

Anche in questo caso, cambiamo l'origine ponendo al posto di x , $x - a$. Si ha:

$$y = y_1 e^{\frac{x^2}{2b_0}}$$

Ponendo ancora $\frac{1}{2b_0} = -\frac{1}{c}$, otteniamo la forma data da Pearson $y = y_0 e^{-\frac{x^2}{c}}$ (3.51).

3.2.4 Le curve secondarie o sotto specie del sistema di Pearson

Le curve secondarie o sotto specie, classificate da Pearson come curve dei tipi VIII, IX, X, XI e XII si ottengono, ad eccezione della curva XII, dalle principali funzioni facendo assumere ai loro parametri determinati valori. Vediamole in ordine.

TIPO VIII

La funzione che rappresenta le curve del tipo VIII, si ottiene quando nella funzione del tipo I

$$va_2 = 0 \text{ e } va_1 < 0.$$

Si ha quindi, ponendo ,

$$y = y_0 \left(1 + \frac{x}{a}\right)^{-m} \quad (3.52)$$

TIPO IX

Anche questa famiglia di curve è un caso particolare del tipo I, che si ottiene quando $va_2 = 0$ e $va_1 > 0$.

Si ha quindi, ponendo $va_1 = m$,

$$y = y_0 \left(1 + \frac{x}{a}\right)^m \quad (3.53)$$

TIPO X

Si tratta di un caso particolare del tipo III, dal quale si ottiene quando $\gamma a = 0$. Pertanto, si ha la funzione:

$$y = y_0 e^{-\gamma x} \quad (3.54)$$

TIPO XI

Si ottiene dalla funzione che rappresenta le curve del tipo VI quando $m_1 = 0$ e $m_2 = m$. In tal caso la funzione assume la forma

$$y = y_0 x^{-m} \quad (3.55)$$

TIPO XII

Il procedimento per ottenere la funzione che rappresenta le curve del tipo XII è più laborioso. Rispetto alle altre curve secondarie, ottenute facendo assumere ai parametri delle funzioni principali determinati valori, in questo caso Pearson agisce sull'equazione differenziale nella forma (3.26) facendo assumere all'espressione $5\beta_2 - 6\beta_1 - 9$ il valore zero. Deducendo dalla relazione indicata il valore $\beta_2 = \frac{1}{5}(6\beta_1 + 9)$, sostituendo a β_2 quest'ultima espressione e ponendo inoltre $\sqrt{\mu_2} = \sigma$ e $\bar{\mu}_2 = \sigma^2$, dalla (3.26) otteniamo l'equazione differenziale:

$$\frac{1}{y} \frac{dy}{dx} = \frac{-2\sqrt{\beta_1}}{\sigma(3 + \beta_1) - \sigma(\sqrt{\beta_1} - \frac{x}{\sigma})^2} \quad (3.56)$$

Risolvendo si ha:

$$\begin{aligned} \log y &= -2\sqrt{\beta_1} \int \frac{1}{-\frac{1}{\sigma}(x^2 - 2\sqrt{\beta_1}\sigma x - 3\sigma^2)} dx \\ \log y &= 2\sigma\sqrt{\beta_1} \int \frac{1}{x^2 - 2\sqrt{\beta_1}\sigma x - 3\sigma^2} dx \end{aligned}$$

Le due radici dell'equazione

$$x^2 - 2\sqrt{\beta_1}\sigma x - 3\sigma^2 = 0$$

Sono:

$$x_1 = \sigma\sqrt{\beta_1} + \sigma\sqrt{3 + \beta_1}$$

$$x_2 = \sigma\sqrt{\beta_1} - \sigma\sqrt{3 + \beta_1}$$

Dunque, si ha:

$$\begin{aligned} \log y &= 2\sigma\sqrt{\beta_1} \int \frac{1}{[x - \sigma(\sqrt{\beta_1} + \sqrt{3 + \beta_1})][x - \sigma(\sqrt{\beta_1} - \sqrt{3 + \beta_1})]} dx \\ \log y &= 2\sigma\sqrt{\beta_1} \int \left\{ \frac{A}{[x - \sigma(\sqrt{\beta_1} + \sqrt{3 + \beta_1})]} + \frac{B}{[x - \sigma(\sqrt{\beta_1} - \sqrt{3 + \beta_1})]} \right\} dx \end{aligned}$$

Applicando il principio d'identità dei polinomi nell'equazione a secondo membro si hanno per A e B i seguenti valori:

$$\begin{cases} A = +\sqrt{\frac{\beta_1}{3 + \beta_1}} \\ B = -\sqrt{\frac{\beta_1}{3 + \beta_1}} \end{cases}$$

Sostituendo i valori ottenuti nell'equazione si ha:

$$\log y = \int \left\{ \frac{\sqrt{\frac{\beta_1}{3+\beta_1}}}{[x - \sigma(\sqrt{\beta_1} + \sqrt{3+\beta_1})]} - \frac{\sqrt{\frac{\beta_1}{3+\beta_1}}}{[x - \sigma(\sqrt{\beta_1} - \sqrt{3+\beta_1})]} \right\} dx$$

$$\log y = \sqrt{\frac{\beta_1}{3+\beta_1}} \left\{ \log [x - \sigma(\sqrt{\beta_1} + \sqrt{3+\beta_1})] - \log [x - \sigma(\sqrt{\beta_1} - \sqrt{3+\beta_1})] \right\} + c$$

$$\log y = \sqrt{\frac{\beta_1}{3+\beta_1}} \left[\log \frac{x - \sigma(\sqrt{\beta_1} + \sqrt{3+\beta_1})}{x - \sigma(\sqrt{\beta_1} - \sqrt{3+\beta_1})} \right] + \log c$$

Togliendo i logaritmi si ha:

$$y = y_0 \left[\frac{x - \sigma(\sqrt{\beta_1} + \sqrt{3+\beta_1})}{x - \sigma(\sqrt{\beta_1} - \sqrt{3+\beta_1})} \right]^{\sqrt{\frac{\beta_1}{3+\beta_1}}} \quad (3.57)$$

Questa funzione, soluzione dell'equazione differenziale (3.56), è diversa da quella che Pearson riporta nel suo articolo (14). La soluzione di Pearson assume la forma:

$$y = y_0 \left[\frac{\sigma(\sqrt{3+\beta_1} + \sqrt{\beta_1}) + x}{\sigma(\sqrt{3+\beta_1} - \sqrt{\beta_1}) - x} \right]^{\sqrt{\frac{\beta_1}{3+\beta_1}}} \quad (3.58)$$

Ebbene è sorprendente rilevare che questa funzione, così come è stata ottenuta da Pearson, non è la soluzione dell'equazione differenziale (3.56). Più precisamente, l'equazione differenziale (3.56) è corretta. Pearson la definisce mediante opportune sostituzioni compiute nell'equazione (3.26), ma la funzione soluzione non corrisponde a quella che ci indica Pearson. Possiamo dimostrare, derivando la funzione (3.58), che l'equazione differenziale di cui è soluzione non è la (3.56), ma una diversa equazione differenziale così definita:

$$\frac{1}{y} \frac{dy}{dx} = \frac{2\sqrt{\beta_1}}{\sigma(3+\beta_1) - \sigma(\sqrt{\beta_1} + \frac{x}{\sigma})^2} \quad (3.59)$$

Non è azzardato affermare che si tratti di un vero e proprio errore matematico. Se si fosse trattato di un errore di stampa, Pearson non avrebbe definito il range di definizione della funzione (3.58) in maniera esatta e proprio corrispondente alla funzione che lui ottiene. Nel fissare l'insieme di definizione si sarebbe accorto dell'errore nella soluzione. Questo ci porta a pensare che la funzione (3.58) è proprio la funzione del tipo XII che Pearson "voleva" ottenere. Ma per ottenere tale funzione avrebbe dovuto modificare l'equazione differenziale (3.56) la quale invece, tenendo conto delle condizioni poste e delle sostituzioni così come indicate da Pearson, è del tutto corretta. Insomma, data la grandezza e competenza di Karl Pearson, ci troviamo di fronte ad un vero e proprio mistero matematico. Ma quello che più sorprende è il fatto che i pochi studiosi che hanno trattato il sistema delle curve di Pearson o non si sono accorti della incongruenza tra la soluzione e l'equazione differenziale, relativamente alle curve del tipo XII del sistema, oppure non se la sono sentita di mettere in discussione il risultato ottenuto da un matematico del calibro di Karl Pearson. Ma in matematica, ed è questa la sua vera grandezza, non vale il principio di autorità, ma l'esattezza dei risultati ottenuti a prescindere da chi li ottiene.

3.2.5 Il metodo di calcolo dei parametri delle funzioni di Pearson

Al pari del grande matematico Joseph Fourier il quale sviluppò, nella "Théorie analytique de la Chaleur" (15), la sua famosa equazione differenziale alle derivate parziali del calore e la risolse mediante il suo innovativo metodo di separazioni delle variabili, facendo largo uso delle serie trigonometriche, le serie di Fourier (16), così anche Karl Pearson, definite le funzioni soluzioni della sua equazione differenziale, con il metodo di interpolazione dei momenti, da lui introdotto nel 1895 (17), riuscì a definirne le formule necessarie per il calcolo dei parametri (18). Gli sviluppi analitici necessari per definire tali formule sono lunghi e laboriosi e non possono trovare spazio in questo articolo. Il lettore interessato può consultare le indicazioni riportate nella nota (18).

4. Commenti e note critiche sul sistema delle curve di Pearson

Il sistema delle curve di Pearson, che il nostro autore espose in tre distinti articoli (19), non tardò ad attirare l'attenzione degli studiosi i quali, nei diversi momenti storici formularono giudizi di apprezzamento, ma anche severe note critiche.

Una combinazione di critica e apprezzamento la possiamo trovare in Benedetto Barberi (20) il quale, riguardo all'equazione differenziale di Pearson, così si esprime: "Una sorgente quasi inesauribile di funzioni di densità per distribuzioni sia simmetriche sia non simmetriche è l'equazione differenziale di Karl Pearson". Ed ancora "è da dire che di tutta la costruzione pearsoniana dal punto di vista operativo, come delle concrete applicazioni, si salvano poche funzioni risultanti da casi molto particolari". Inoltre, prosegue Barberi "l'inconveniente delle funzioni dell'autore in questione è il grande numero di parametri arbitrari e di nessun significato dal punto di vista dell'interpretazione delle leggi di distribuzione" (21).

Queste severe osservazioni sono mitigate da alcune considerazioni riguardanti l'applicazione di alcune distribuzioni della famiglia di Pearson agli organismi tipo e in particolare agli organismi economici (22). "Una di queste distribuzioni particolarmente adatte a rappresentare le caratteristiche strutturali degli organismi economici, conclude Barberi, è la distribuzione analiticamente specificata dalla funzione di densità del tipo V nella sistematica pearsoniana" (23). Di grande significato è quanto afferma Gustavo Barbensi (24): "Le distribuzioni che si presentano nel campo biologico offrono, per lo più, la caratteristica di iniziare con frequenze minime che poi raggiungono un valore massimo per ritornare ad essere nuovamente minime. Si tratta dunque per lo più di curve più o meno asimmetriche, eccezionalmente simmetriche" (25). E ancora: "La concezione del Quetelet che le variazioni dei fenomeni biologici siano in dipendenza del caso e che possono quindi esprimersi mediante la legge degli errori, se ha suscitato discussioni quanto alla interpretazione ha sempre trovato seguaci quando si è trattato di trovare una rappresentazione analitica o geometrica del fenomeno. Si ricollega a questo fatto e ne è in certo qual modo una conferma il grande successo del sistema delle curve di Pearson, al quale i biologi sogliono molto volentieri ricorrere quando vogliano procedere alla interpolazione delle distribuzioni. Afferma infatti lo Stefensen che quel successo dipende da che esse debbano la loro pratica utilità all'essere derivate dalla sorgente naturale del calcolo delle probabilità" (26).

Se per Barberi, come abbiamo visto, i parametri arbitrari delle funzioni di Pearson rappresentano un inconveniente, non lo sono invece per Pierpaolo Luzzatto Fegiz (27) il quale, riguardo alle curve di Pearson, così si esprime: "Tali curve traggono la loro importanza dal fatto che, contenendo quattro costanti arbitrarie, oltre a quella che deriva dall'integrazione, si piegano a rappresentare con successo le diverse forme di curve di frequenza quali empiricamente si manifestano nella biometria" (28). Parole di apprezzamento sono anche quelle di Francesco Brambilla (29) il quale, riferendosi alle soluzioni dell'equazione di Pearson, sostiene che esse "formano un sistema di riferimento fondamentale noto col nome di Sistema di Karl Pearson" (30). Ed ancora: "Il sistema di riferimento di Karl Pearson ha posto in evidenza la possibilità di avere delle famiglie di curve di riferimento utilizzando i primi quattro momenti e di possedere un criterio (l'indice K) atto alla identificazione rapida del tipo di famiglia cui una variabile osservata può appartenere" (31). Italo Scardovi (32) sottolinea che: "l'elemento di novità introdotto dalle formule di Pearson sta nel preferire alla ricerca di una funzione analitica la determinazione e la soluzione della sua equazione differenziale, essendo spesso le equazioni differenziali più semplici e più significative delle funzioni corrispondenti" (33).

Un'analisi più approfondita del sistema di Pearson viene svolta da Felice Vinci (34), il quale distingue l'aspetto probabilistico all'origine dell'equazione differenziale di Pearson da quello che considera tale equazione come il risultato della traduzione differenziale della dinamica delle distribuzioni di frequenza. Vinci critica fortemente la deduzione dell'equazione di Pearson della distribuzione ipergeometrica. Per Vinci "anche nello schema binomiale di Bernoulli avremmo potuto dedurre diversi tipi di curve" (35). Dunque, nulla di nuovo egli afferma. Ed ancora: "Diverso è il giudizio se si rinuncia a far discendere le sue curve dalla distribuzione ipergeometrica" (36).

Semplici considerazioni attinenti al comportamento delle distribuzioni di frequenze, sottolinea Vinci, "permettono di partire direttamente dall'equazione di Pearson e quindi ottenere, oltre alla curva gaussiana, quei diversi tipi di curve, l'utilità delle quali potrebbe riconoscersi nella loro attitudine a rispecchiare molte forme di distribuzione e nell'ausilio che esse potrebbero dare all'applicazione di qualche speciale schema teorico" (37). Tuttavia, sottolinea ancora Vinci "è necessario avvertire riguardo all'arbitrarietà dell'ipotesi che la funzione e denominatore dell'equazione differenziale di Pearson sia rappresentabile con il trinomio" (38). Da ultimo anche Marcello Boldrini (39), uno dei maggiori statistici italiani della prima metà del secolo scorso, non fece mancare le sue osservazioni critiche al sistema di Pearson. In particolare, così si esprime nel suo testo sull'interpolazione statistica: "All'equazione differenziale di Pearson si riconosce interesse solo come formula empirica, in quanto la sua integrazione conduce a funzioni interpolatrici atte a rappresentare una larga categoria di distribuzioni di frequenza; ma a queste funzioni teoriche di frequenza si attribuisce significato puramente descrittivo e non interpretativo, perché non sono in grado di illuminarci sulle circostanze influenti che determinano il comportamento della ripartizione dei caratteri considerati" (40).

5. Osservazioni conclusive

Dunque, come abbiamo visto, diversi statistici, tra i più autorevoli della prima metà del '900, pur non facendo mancare severe critiche nel complesso hanno accolto favorevolmente la teoria delle curve di Pearson. Ognuno ha saputo cogliere qualche aspetto della sofisticata costruzione teorica. Ora, a prescindere dalle loro osservazioni, possiamo dire che sicuramente la caratteristica più importante e positiva del sistema delle curve di Pearson è la propria capacità di esprimere una particolare modalità di unificazione di funzioni diverse, in quanto nonostante la loro diversità sono tutte riconducibili, previa definizione iniziale dei parametri, alla stessa equazione differenziale. Ciò consente di inquadrare tale sistema nella categoria dei modelli matematici. Considerare le diverse funzioni di Pearson dei modelli matematici comporta un ridimensionamento della critica del grande statistico Marcello Boldrini e questo perché un modello matematico, citando il matematico Giorgio Israel è "uno schema vuoto di realtà possibili" (41). Ora, da questa definizione consegue l'elevata astrazione di un modello matematico, la sua generalità e potenzialità (42). Queste caratteristiche le ritroviamo ampiamente nelle funzioni di Pearson. I modelli matematici non pretendono di stabilire un rapporto biunivoco con il fenomeno da studiare. Modelli diversi possono descrivere lo stesso fenomeno e fenomeni diversi possono essere descritti dallo stesso modello matematico. E questo è proprio ciò che avviene con le funzioni di Pearson. Esse costituiscono una sorta di repertorio teorico dal quale attingere per descrivere e studiare i fenomeni che di volta in volta ci interessano. La capacità delle funzioni di Pearson di descrivere un fenomeno senza la pretesa di volerlo spiegare ci consente di interpretare in senso positivo la critica di Marcello Boldrini il quale, come abbiamo visto, attribuiva alle funzioni di Pearson una capacità descrittiva, ma non interpretativa. Ma è necessario precisare che Boldrini con l'espressione "capacità

descrittiva" non intendeva evidenziare un aspetto positivo delle funzioni di Pearson, ma un loro limite teorico, riguardante la possibilità di spiegare e interpretare i fenomeni, ma la crescente importanza dei modelli matematici nella descrizione e nello studio dei più svariati fenomeni della vita reale riesce a rendere più che mai attuali e versatili le funzioni di Karl Pearson (43).

Note

- [1] K. PEARSON, Mathematical Contribution to the Theory of evolutions, IV: *On the probable errors of frequency constants and on the Influence of random selection on variation and correlation*, in "Philosophical Transactions of the Royal Society", Serie A, 191, 1898, pag 229-311.
- [2] K. PEARSON, *On the criterion that a given System of deviations from the probable in the case of a correlated System of variables is such that it can be reasonably supposed to Have Arisen from random sampling* Phil. mag., 6, p. 157-175, 1900.
- [3] Più precisamente, Pearson riteneva che ci fosse "un considerevole pericolo nell'applicare i metodi delle scienze esatte ai problemi delle discipline descrittive, sia che si tratti di problemi di ereditarietà che di politica economica" NORTON B.J., Karl Pearson and Statistics: *The Social Origins of Scientific Innovation*, in "Social Studies of Science", 8, 1, 3-34.
- [4] Tipo di granchio.
- [5] F. GALTON, *Natural Inheritance*, New York, 1889.
- [6] J. I. PIOVANI, *Alle origini della statistica moderna*, F. Angeli, 2006, pag. 52;
- [7] K. PEARSON, *The Grammar of Science*, Cambridge University Press, 1892;
- [8] C. M. MADRID CASADO, FISHER, *L'inferenza statistica, Probabilmente sì, probabilmente no*, RBA, 2014, pag. 49.
- [9] Uno dei maggiori contributi di Kuno Fisher alla filosofia fu la distinzione tra empirismo e razionalismo. Nell'empirismo la conoscenza, sosteneva Kuno Fisher, può essere acquisita attraverso l'esperienza, mentre nel razionalismo l'acquisizione avviene tramite categorie e concetti. Riguardo ad Ernst Mach, uno dei più influenti filosofi della scienza del periodo a cavallo tra l'800 e il '900, la conoscenza non può prescindere dai nostri sensi. Pertanto, lottò, da autentico positivista per l'eliminazione dei concetti metafisici ancora presenti nelle teorie fisiche. In particolare, criticò la nozione di spazio assoluto e in questo si può considerare un precursore della teoria della relatività ristretta.
- [10] Per una esposizione più dettagliata si veda K. PEARSON, *The Grammar of Science*, cit., pagg. 92-135.

- [11] C. M. MADRID CASADO, FISHER, cit., pag.50.
- [12] A. MAROS DELL'ORO, *Storia del metodo statistico*, Giuffrè, 1976, pag. 102.
- [13] K. PEARSON, *Contributions to the Mathematical Theory of Evolution - II. Skew variation in homogeneous material*. Philos. Trans. R. Soc. A 186, 343-414, 1895.
- K. PEARSON, *Contributions to the Mathematical Theory of Evolution - X. Supplement to the Memoir on Skew variation*. Philos. Trans. R. Soc. A 197, 443 - 459, 1901.
- K. PEARSON, *Contributions to the Mathematical Theory of Evolution - XIX. Second Supplement to a Memoir on Skew Variation*. Philos. Trans. R. Soc. A 216, 429 - 457, 1916.
- [14] Si vedano le pagg. 437 e seguenti del terzo articolo di K. Pearson del 1916.
- [15] La Thèorie analytique de la chaleur venne pubblicata nel 1822. Essa costituisce l'estensione e il completamento di una precedente memoria che Fourier presentò all'Istituto di Francia nel 1807, rimasta manoscritta e pubblicata nel 1972. Si vedano "GRATTAN-GUINNES I, RAVETZ J.R., JOSEPH FOURIER, 1768-1830, Cambridge, 1972 e " U. BOTTAZZINI, *il calcolo sublime: Storia dell'analisi infinitesimale da Eulero a Weierstrass*, 1981, pp. 58.e 59.
- [16] U. BOTTAZZINI, op. cit., pag 59 e A.N. TICHONOV, A.A. SAMARSKI, *Equazioni della fisica matematica, traduzione italiana*, mir, 1981, pag 86.
- [17] Sulle origini del metodo dei momenti si veda L. Vajani, *Statistica descrittiva*, ETAS, 1990, pag. 431. K. Pearson espone tale metodo nell'articolo "Skew variations in homogeneous material" in *Phil. Trans. of the Royal Society*, 1895.
- [18] Il metodo dei momenti si può trovare esposto in maniera dettagliata in tutti i testi di statistica. Per il calcolo dei parametri delle curve di Pearson si possono consultare i testi:
- W. Palin Elderton, *Frequency Curves and Correlation*, Cambridge University Press, 1906.
 - A. Costanzo, *Statistica*, Giuffrè, 1973, appendice matematica, p. IV, pagg. 313-353.
- [19] Si veda la nota 13.
- [20] Benedetto Barberi, nato a Cittareale nel 1901 e morto a Roma nel 1976, è stato un importante statistico italiano. Direttore Generale dell'ISTAT dal 1946 al 1963, anno nel quale divenne professore di statistica economica nella facoltà di Scienze Statistiche dell'Università di Roma "La Sapienza".
- [21] B. Barberi, *Statistica, Teoria ed applicazioni*, Edizioni CERES, Roma, 1971, pagg. 163-164.

- [22] B. Barberi definisce gli organismi tipo delle realtà distinte dagli individui empirici che pro-tempore le configurano.
- [23] B. Barberi, op. cit., pag. 167.
- [24] Gustavo Barbensi, nato a Roma il 21/09/1875 e morto il 07/07/1974, è stato un medico e statistico noto, autore di un importante testo di statistica biometrica citato nelle note successive.
- [25] G. Barbensi, *Elementi di Metodologia Biometrica*, Editore Luigi Niccolai, 1940, pag. 210.
- [26] G. Barbensi, op. cit., pag. 211.
- [27] Pierpaolo Luzzatto Fegiz, statistico italiano, nato a Trieste nel 1900 ed ivi morto nel 1989, è stato professore nelle Università di Trieste dal 1931 e in quella di Roma dal 1961. Fondatore nel 1946 e direttore dell'Istituto Doxa, divenne successivamente nel 1973 socio dell'Accademia Nazionale dei Lincei.
- [28] P. Luzzatto Fegiz, *Statistica demografica ed economica*, UTET, 1962, pag. 643.
- [29] Francesco Brambilla, nato a Milano nel 1913 ed ivi morto nel 1996, è stato un importante statistico italiano. Professore di Statistica nell'Università Statale di Milano, nella quale ha fondato e diretto il centro per la Ricerca Operativa. Tra le sue opere più importanti è da menzionare il *Trattato di Statistica e Ricerca Operativa* in 5 volumi pubblicato tra il 1968 e il 1970.
- [30] F. Brambilla, *Trattato di Statistica e Ricerca Operativa vol. I, La Variabilità*, UTET, 1968, pag. 105.
- [31] F. Brambilla, op. cit., pag. 111.
- [32] Italo Scardovi, Accademico dei Lincei, nato a Bologna il 13/12/1928, è professore Emerito di Teoria dell'Inferenza Statistica presso l'Università degli Studi di Bologna.
- [33] I. Scardovi, *Argomenti di Metodologia Statistica*, Patron, Editore Bologna, 1975, pag. 360.
- [34] Statistico ed Economista, Felice Vinci nasce a Palermo nel 1890 e muore a Roma nel 1962. Professore dal 1942 prima di Statistica e poi di Economia Politica nell'Università Statale di Milano. Fondatore della "Rivista Italiana di Statistica", è stato tra i primi statistici italiani a impostare la metodologia statistica su basi strettamente logico-matematiche.
- [35] F. Vinci, *Manuale di Statistica, volume secondo*, Zanichelli Editore, 1934, pag. 92.
- [36] F. Vinci, op. cit., pag. 93.
- [37] F. Vinci, op. cit., pag. 93.
- [38] F. Vinci, op. cit., pag. 93.

- [39] Marcello Boldrini, uno dei più noti statistici italiani, è nato a Matelica nel 1890 e morto a Milano nel 1969. Professore dal 1922 di statistica e demografia nelle Università di Messina, Padova, Milano-Bicocca e Cattolica a Roma. È stato socio dell'Accademia dei Lincei e dell'Accademia Pontificia. Successivamente, venne nominato Presidente dell'AGIP e negli anni 1962-67, vicepresidente e poi presidente dell'ENI.
- [40] M. Boldrini, A. Uggè, *L'interpolazione statistica*, Giuffrè, Milano, 1950, pag. 230.
- [41] G. Israel, A. Millan Gasca, *Pensare in matematica*, Zanichelli Editore, 2012, pag. 435.
- [42] La letteratura sui modelli matematici è molto vasta. Un testo chiaro e abbastanza esauriente è G. Israel, *La visione matematica della realtà*, Bari, 1996.
- [43] Vi sono diversi e importanti fenomeni che possono essere descritti dalle funzioni di Pearson. Ad esempio, nell'ambito dell'ingegneria civile e dei trasporti possiamo citare i modelli matematici del flusso del traffico e quelli utilizzati negli studi di idrologia. Al riguardo si possono consultare le seguenti fonti:
- a) M. Ahmed Rehmetullah, *Toward Modelling for Crashes. A Low-Cost Adaptive Methodology for Karachi*, World Academy of Science, Engineering and Technology, 48, 2008, pag. 245.
 - b) *Distribuzioni di Pearson, applicazioni* in "Encyclopedia, Science, News e Research Reviews".
 - c) M. Roche, *Hydrologie de Surface*, Gauthier-Villars, Paris, 1963.
 - d) Diversi articoli sulle applicazioni in idrologia sono reperibili all'indirizzo di posta elettronica idrologia@polito, del Politecnico di Torino.
 - e) Le funzioni di Pearson non si applicano soltanto in alcuni settori dell'ingegneria civile. Queste distribuzioni, e in particolar modo quelle del tipo III, vengono usate anche nell'analisi dei mercati finanziari, in quanto consentono una possibilità di parametrizzazione più simile al modo di pensare di chi opera su tali mercati. Tra le diverse variabili casuali usate correntemente per descrivere la natura stocastica della volatilità dei cambi e delle azioni, le variabili di Pearson sono tra le più impiegate. Si veda al riguardo la fonte già citata (b). Infine, altri casi di applicazione delle funzioni di Pearson sono riportati nel testo di William Palin Elderton già citato nella nota 18.

Le "infinite" identità del fattoriale, dai reali agli ipercomplessi

The "infinite" identities of the factorial: from real to hypercomplex numbers

Giovanni Di Mare¹ e Domenico Veca²

Abstract

Lo studio espone la scoperta di una possibile nuova identità, ottenuta dal aver considerato le differenze iterate delle potenze n -esime di $n + 1$ numeri naturali consecutivi. Tramite una rappresentazione ad albero delle differenze, siamo giunti al valore di n -fattoriale. La seguente identità è stata provata nell'insieme dei numeri reali, la cui dimostrazione ci ha portato all'enunciazione di due nuove possibili proprietà del triangolo di Tartaglia-Pascal, una delle quali generalizza la nota proprietà della somma, a segni alterni, dei numeri di una qualsiasi riga che risulta uguale a 0. Abbiamo trovato delle estensioni dell'identità nei numeri complessi, nei quaternioni, negli ottetti e nei sedenioni, con i polinomi di Bernstein e con l'insieme di Mandelbrot. Inoltre nell'articolo vi sono tutte le prove computazionali in linguaggio Python.

Introduzione

Iniziamo a ragionare con il caso banale $n = 2$. Consideriamo le potenze alla seconda dei 3 numeri consecutivi 2, 3 e 4 e poniamo $x_1 = 4$, $x_2 = 9$ e $x_3 = 16$.

Osserviamo che $x_3 - x_2 = 7$ e che $x_2 - x_1 = 5$; se facciamo la differenza delle due differenze precedenti otteniamo:

$$(x_3 - x_2) - (x_2 - x_1) = x_3 - 2x_2 + x_1 = 2$$

Invece se consideriamo tre qualunque numeri reali consecutivi $x - 2$, $x - 1$ e x e allora l'equazione assume la forma:

$$x^2 - 2(x - 1)^2 + (x - 2)^2 = 2 \tag{1}$$

¹domenicoveca5@gmail.com

²giova257@icloud.com

L'equazione è sempre vera per ogni $x \in \mathbb{R}$. Per verificarla basta sviluppare i quadrati del primo membro e semplificando si ottiene:

$$x^2 - 2x^2 + 4x - 2 + x^2 - 4x + 4 = 2$$

Dall'equazione notiamo che i coefficienti della somma dei quadrati sono 1, -2, 1 e il secondo membro è uguale a 2. Passiamo al caso $n = 3$, quindi consideriamo i cubi di 4 dei numeri consecutivi 2, 3, 4 e 5: (vedremo che per lo stesso motivo la scelta dei numeri consecutivi è arbitraria) 8, 27, 64 e 125. Poniamo $x_1 = 8$, $x_2 = 27$, $x_3 = 64$, $x_4 = 125$. Se calcoliamo le differenze $x_{i+1} - x_i$ per $i = 1, 2, 3$ otteniamo: $x_2 - x_1 = 19$, $x_3 - x_2 = 37$, $x_4 - x_3 = 61$. Come nel caso precedente, calcoliamo le differenze: $(x_{i+2} - x_{i+1}) - (x_{i+1} - x_i)$ ovvero $(x_4 - x_3) - (x_3 - x_2) = 24$ e $(x_3 - x_2) - (x_2 - x_1) = 18$ e infine la differenza:

$$[(x_4 - x_3) - (x_3 - x_2)] - [(x_3 - x_2) - (x_2 - x_1)] = x_4 - 3x_3 + 3x_2 - x_1 = 6$$

Allora se consideriamo quattro qualunque numeri reali consecutivi $x - 3$, $x - 2$, $x - 1$ e x l'equazione diventa:

$$x^3 - 3(x - 1)^3 + 3(x - 2)^3 - (x - 3)^3 = 6 \quad (2)$$

Osserviamo che l'equazione, come la (1) è un'identità in $\mathbb{R}$. I coefficienti della somma dei cubi sono: 1, -3, 3, -1 e il secondo membro è uguale a 6. Se ripetiamo il procedimento, per il caso $n = 4$, l'equazione diventa:

$$x^4 - 4(x - 1)^4 + 6(x - 2)^4 - 4(x - 3)^4 + (x - 4)^4 = 24 \quad (3)$$

dove i coefficienti della somma delle potenze alla quarta sono 1, -4, 6, 4 e 1 e il secondo membro è 24. La (4) è anch'essa un'identità in $\mathbb{R}$.

Per il salto della fiducia i coefficienti nel caso n sono i numeri, con segno alternato, della n -esima riga del triangolo di Tartaglia - Pascal e il termine noto del secondo membro è uguale al valore di $n!$. Infatti $2! = 2$, $3! = 6$ e $4! = 24$.

Identità in $\mathbb{R}$

Proprietà del triangolo di Tartaglia-Pascal

Per dimostrare l'identità nel caso generale n bisogna introdurre due proprietà del triangolo di Tartaglia-Pascal, la prima più nota, mentre la seconda è stata trovata mentre si cercava la dimostrazione, studiando i casi più semplici, come $n = 2, 3, 4$, ecc.

Proprietà 1. La somma a segni alterni di tutti i numeri del n -esima riga, per $n \geq 1$ è uguale a 0.

$$\sum_{k=0}^n (-1)^k \binom{n}{k} = 0. \quad (4)$$

Proprietà 2. La somma a segni alterni dei prodotti dei coefficienti binomiali n su k e delle potenze i -esime dei numeri $-k$, per k da 1 a n , è uguale a:

$$\sum_{k=1}^n (-1)^k \binom{n}{k} (-k)^i = \begin{cases} 0 & \text{per } i = 1, \dots, n-1 \\ n! & \text{per } i = n \end{cases} \quad (5)$$

Per esempio, se consideriamo tutti i numeri della riga $n = 5$ del triangolo di Tartaglia, escluso il primo $\binom{5}{0} = 1$ si ha:

- per $i = 1$: $5 - 20 + 30 - 20 + 5 = 0$;
- per $i = 2$: $-5 + 40 - 90 + 80 - 25 = 0$;
- per $i = 3$: $5 - 80 + 270 - 320 + 125 = 0$
- per $i = 4$: $-5 + 160 - 810 + 1280 - 625 = 0$
- per $i = 5$: $5 - 320 + 2430 - 5120 + 3125 = 120 = 5!$

Teorema dell'identità di n-fattoriale in $\mathbb{R}$

Adesso possiamo dimostrare il caso reale per un n naturale.

Teorema 1. Sia $x \in \mathbb{R}$, allora vale la seguente identità:

$$\sum_{k=0}^n (-1)^k \binom{n}{k} (x - k)^n = n! \quad (6)$$

dove $n \in \mathbb{N}$, con $n \geq 1$ e $\binom{n}{k}$ è il coefficiente binomiale di Newton. *Dimostrazione.* Con lo sviluppo del binomio di Newton, scriviamo il polinomio

$$(x - k)^n = \sum_{i=0}^n \binom{n}{i} x^{n-i} (-k)^i.$$

Allora la (6), esplicitando x^n , diventa:

$$\begin{aligned} x^n + \sum_{k=1}^n (-1)^k \binom{n}{k} \left[\sum_{i=0}^n \binom{n}{i} x^{n-i} (-k)^i \right] &= x^n + \sum_{i=0}^n \binom{n}{i} \left[\sum_{k=1}^n (-1)^k \binom{n}{k} x^{n-i} (-k)^i \right] = \\ &= \left[1 + \sum_{k=1}^n (-1)^k \binom{n}{k} (-k)^0 \right] x^n + \sum_{i=1}^{n-1} \binom{n}{i} \left[\sum_{k=1}^n (-1)^k \binom{n}{k} (-k)^i \right] x^{n-i} + \\ &+ \binom{n}{n} \left[\sum_{k=1}^n (-1)^k \binom{n}{k} (-k)^n \right] x^0. \end{aligned}$$

Poichè $1 + \sum_{k=1}^n (-1)^k \binom{n}{k} (-k)^0$ si può riscrivere $\sum_{k=0}^n (-1)^k \binom{n}{k}$ allora, per la (4), quest'ultima è uguale a 0. Inoltre, per la (5), $\sum_{k=1}^n (-1)^k \binom{n}{k} (-k)^i = 0$, per $i = 1, \dots, n-1$ e $\sum_{k=1}^n (-1)^k \binom{n}{k} (-k)^n = n!$ Pertanto risulta:

$$0 + 0 + n! = n!$$

□

Identità in $\mathbb{C}$

I nostri studi ci hanno spinto a porre la seguente domanda: è possibile la seguente identità valga anche per i numeri complessi? Se sì, per quali casi?

Dato un numero complesso $x = a + bi$ esponiamo i seguenti casi:

- 1) Se fissiamo la parte immaginaria b e decrementiamo la parte reale a per k volte, l'identità (6) risulta valida. Le prove effettuate sono state fatte a mano per casi $n = 2, 3, 4$, e per via computazionale (vedi "Parte computazionale" 8.2).
- 2) Se decrementiamo la parte immaginaria b per k volte e fissiamo la parte reale a , la formula assume la seguente forma:

$$\sum_{k=0}^n (-1)^k \binom{n}{k} [a + (b - k)i]^n = n! \cdot i^n \quad (7)$$

- 3) Se decrementiamo per k volte sia la parte reale che la parte immaginaria, otteniamo 4 identità che, con un ciclo di 4, si ripetono all'infinito. Infatti la sommatoria

$$\sum_{k=0}^n (-1)^k \binom{n}{k} [(a - k) + (b - k)i]^n$$

è uguale a, al variare di $m \in \mathbb{N}$:

$$\begin{aligned} &\text{- per } n = 4m, \\ &\quad (-1)^m 2^{\frac{n}{2}} n! (1 + 0i) \end{aligned} \quad (8)$$

$$\begin{aligned} &\text{- per } n = 4m + 1, \\ &\quad (-1)^m 2^{\frac{n-1}{2}} n! (1 + i) \end{aligned} \quad (9)$$

$$\begin{aligned} &\text{- per } n = 4m + 2, \\ &\quad (-1)^m 2^{\frac{n}{2}} n! (0 + i) \end{aligned} \quad (10)$$

$$\begin{aligned} &\text{- per } n = 4m + 3, \\ &\quad (-1)^m 2^{\frac{n-1}{2}} n! (-1 + i) \end{aligned} \quad (11)$$

Osservazione 1

Se volessimo rappresentare i secondi membri delle 4 identità sul piano Argan-Gauss otterremo infiniti vettori, con punto di applicazione nell'origine degli assi, multipli degli 4 vettori: $(1, 0)$, $(1, 1)$, $(0, 1)$ e $(-1, 1)$.

Identità nei quaternioni

Esponiamo i seguenti casi:

- 1) L'identità (6), applicata ai quaternioni, è valida quando si decrementa solamente la parte reale.
- 2) L'identità (7), applicata ai quaternioni decrementando una sola qualsiasi unità immaginaria (i, j e k), risulta valida. Le identità (8), (9), (10) e (11) sono vere se vengono decrementate la parte reale e una delle parti immaginarie.
- 3) Invece se decrementiamo due qualsiasi parti immaginarie, per esempio scegliamo la i e la k la sommatoria

$$\sum_{l=0}^n (-1)^l \binom{n}{l} [a + (b-l)i + cj + (d-l)k]^n$$

al variare di $m \in \mathbb{N}$, è uguale a:

- per $n = 2m$,

$$(-1)^m 2^{\frac{n}{2}} n! \quad (12)$$

- per $n = 2m + 1$,

$$(-1)^m 2^{\frac{n-1}{2}} n! (i + k) \quad (13)$$

- 4) Per il decremento della parte reale e le due parti immaginarie qualsiasi, nel nostro caso (reale, i, j) non sono note le identità. Siamo lieti di invitare il lettore a trovarle.
- 5) Nel caso del decremento delle 3 parti immaginarie la sommatoria è uguale a, al variare di $m \in \mathbb{N}$:

- per $n=2m$,

$$(-1)^m 3^{\frac{n}{2}} n! \quad (14)$$

- per $n=2m+1$,

$$(-1)^m 3^{\frac{n-1}{2}} n! (i + j + k). \quad (15)$$

Osservazione 2

Se consideriamo le uguaglianze (12) e (14) e le (13) e (15) sono simili, infatti esse differiscono per la base delle potenze, che cresce da 2 a 3, e inoltre viene aggiunta la terza componente immaginaria.

6) Nel caso del decremento di tutte le parti, sia reale che immaginarie, la sommatoria è uguale a, con $m \in \mathbb{N}$:

$$\begin{aligned} & \text{- per } n = 3m, \\ & \qquad \qquad \qquad 2^n n! \end{aligned} \tag{16}$$

$$\begin{aligned} & \text{- per } n = 3m + 1, \\ & \qquad \qquad \qquad (-1)^m 2^{n-1} n! (1 + i + j + k) \end{aligned} \tag{17}$$

$$\begin{aligned} & \text{- per } n = 3m + 2, \\ & \qquad \qquad \qquad (-1)^{m+1} 2^{n-1} n! (-1 + i + j + k) \end{aligned} \tag{18}$$

Identità negli ipercomplessi

Congetture:

1) Dai dati sperimentali, tutte le identità trovate posso essere estese anche agli ottetti, senioni, ecc., tranne quelle in cui decrementiamo sia la parte reale e le parti immaginarie. Se ci restringiamo al solo decremento delle parti immaginarie, dall' **osservazione 2** e per il salto della fiducia, si può enunciare la seguente:

Regola generale

Le identità con il decremento da 2 fino a $2^n - 1$ parti immaginarie di un qualunque ipercomplesso, di dimensione 2^n , si ricavano partendo dalle identità (12) e (13); dove le basi delle potenze che hanno come esponenti $\frac{n}{2}$ (per la (12)) e $\frac{n-1}{2}$ (per la (13)) sono uguali al numero delle parti immaginarie decrementate. L'ultimo fattore moltiplicativo sarà uguale alla somma delle unità immaginarie interessate.

2) Inoltre, come da studi sperimentali, si ritiene che l'identità (6) per i reali possa essere valida, sostituendo al posto della x , un qualsiasi polinomio di Bernstein e un qualsiasi numero di un insieme di Mandelbrot.

Riguardo le congetture esposte, gli autori si riservano di approfondirle eventualmente in futuro.

Altre proprietà del triangolo di Tartaglia-Pascal

Riprendiamo in esame l'esempio dei casi i , per $n = 5$, della **proprietà 2**. Se raccogliamo e dividiamo entrambi i membri per 5 nei casi i dispari e -5 nei casi i pari otteniamo:

- per $i = 1$:

$$1 - 4 + 6 - 4 + 1 = 0$$

- per $i = 2$:

$$1 - 4 \cdot 2 + 6 \cdot 3 - 4 \cdot 4 + 1 \cdot 5 = 0$$

- per $i = 3$:

$$1 - 4 \cdot 2^2 - 6 \cdot 3^2 + 4 \cdot 4^2 - 1 \cdot 5^2 = 0$$

- per $i = 4$:

$$1 - 4 \cdot 2^3 - 6 \cdot 3^3 + 4 \cdot 4^3 - 1 \cdot 5^3 = 0$$

- per $i = 5$:

$$1 - 4 \cdot 2^4 - 6 \cdot 3^4 + 4 \cdot 4^4 - 1 \cdot 5^4 = 24 = 4!$$

Dai risultati ottenuti possiamo formulare la seguente proprietà:

Proprietà 3. La somma a segni alterni dei prodotti dei coefficienti binomiali n su k e delle potenze i - esime dei $k + 1$ numeri consecutivi, al variare di k da 0 a n , è uguale a:

$$\sum_{k=0}^n (-1)^k \binom{n}{k} (k+1)^i = \begin{cases} 0 & \text{per } i = 0, \dots, n-1 \\ n! & \text{per } i = n \end{cases} \quad (19)$$

Notiamo che la **proprietà 1** corrisponde al caso $i = 0$. Abbiamo provato a rispondere alla seguente domanda: la sommatoria come si comporta se i numeri $k + 1$ hanno esponenti interi negativi? Per $i = -1$ siamo riusciti a trovare la seguente uguaglianza:

Proprietà 4. La somma a segni alterni dei prodotti dei coefficienti binomiali n su k e delle potenze dei numeri $(k + 1)^{-1}$, al variare di k da 0 a n , è uguale a:

$$\sum_{k=0}^n (-1)^k \binom{n}{k} \frac{1}{k+1} = \frac{1}{n+1} \quad (20)$$

Parte computazionale

In questa sezione abbiamo scritto le identità in linguaggio Python per le prove empirico-computazioni di verifica e di riformalizzazione delle stesse.

x appartenente all'insieme dei reali

```
mp.dps = 19999
x = mpc(float(input("Inserisci il valore per x: ")))
n = (int(input("Inserisci il valore per n: ")))
nd = list(range(0, n + 1))
k = 0
result = 0.0
```

```

for k in nd:
first = mpc(math.pow(-1, k))
comb = mpc(math.comb(n, k))
second = mpc((x - k) ** n)
prod = mpc(first * comb * second)
result = mpc(result) + prod
factorial = mpc(math.factorial(n))
difference = result - factorial
difference = mpc(difference)
print(difference)

```

Descrizione

Tale codice, dove la x rappresenta un numero reale, fornisce all'utilizzatore sempre un'identità sotto le dovute condizioni di calcolo imposte dal limite di accuratezza dei risultati di mp.dps. Quindi, a titolo d'esempio, il codice così impostato fornirà sempre un'identità per valori della n fino a 100 e valori arbitrari della x , se calcolabili accuratamente con mp.dps.

x appartenente all'insieme dei numeri complessi (con decremento della sola parte reale)

```

mp.dps = 1999999
n = (int(input("Inserisci il valore per n: ")))
xpartereale = mpc(float(input("Inserisci la parte reale del numero complesso ")))
xparteimmaginaria = (float(input("Inserisci la parte immaginaria del numero
complesso ")))
x = mpc(complex(xpartereale, xparteimmaginaria))
nd = list(range(0, n + 1))
k = 0
result = 0.0
for k in nd:
first = mpc(math.pow(-1, k))
comb = mpc(math.comb(n, k))
second = mpc((x - k) ** n)
prod = mpc(first * comb * second)
result = mpc(result + prod)
factorial = math.factorial(n)
difference = mpc(result - factorial)
print(difference)

```

Descrizione

Tale codice, dove la x rappresenta un numero complesso, fornisce all'utente sempre un'identità sotto le dovute condizioni di calcolo imposte dal limite di accuratezza dei risultati di mp.dps. Quindi, a titolo d'esempio, il codice così impostato fornirà sempre un'identità per valori della n fino a 100 e valori arbitrari della x , se calcolabili accuratamente con mp.dps. In questo caso a essere decrementata è la sola parte reale del numero complesso.

 x appartenente all'insieme dei numeri complessi (con decremento della sola parte immaginaria)

```
mp.dps = 1999999
n = 2
x = mpc(complex(3, -10.3))
nd = list(range(0, n + 1))
k = 0
result = 0.0
for k in nd:
    first = mpc(math.pow(-1, k))
    comb = mpc(math.comb(n, k))
    second = mpc(complex(7, 3 - k) ** n)
    prod = mpc(first * comb * second)
    result = mpc(result + prod)
print(result)
```

Descrizione

In questo caso il risultato è $2! \cdot i^2 = -2$.

 x appartenente all'insieme dei numeri complessi (con decremento sia della parte reale sia della parte immaginaria)

```
n = 10
nd = list(range(0, n + 1))
k = 0
result = 0.00
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    x = complex(8 + 11j)
    kk = complex(8 - k, 11 - k)
    last = kk ** n
    prod = first * comb * last
    result = result + prod
print(result)
```

```
# n = 0: 1+0j
# n = 1: 1+1j
# n = 2: 4j
# n = 3: -12+12j
# n = 4: -96+0j
# n = 5: -480-480j
# n = 6: 5760j
# n = 7: 40320-40320j
# n = 8: 645120+0j
# n = 9: 5806080+5806080j
```

x appartenente all'insieme dei quaternioni (con decremento della sola parte reale)

```
n = 9
nd = list(range(0, n + 1))
x = Quaternion(2,7,3,2)
k = 0
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    quaternionminusk = (x - k) ** n
    prod = first * comb * quaternionminusk
    result = result + prod
factorial = math.factorial(n)
difference = result - factorial
print(f'result:result; factorial: factorial; difference:difference')
result:+362880.000 -0.000i -0.000j -0.000k;
factorial: 362880;
difference:+0.000 -0.000i -0.000j -0.000k
```

Descrizione

In tale codice, dove la x rappresenta un quaternione, si è scelto arbitrariamente $n = 9$ e un quaternione con coefficienti 2, 7, 3 e 2. La "difference" corrisponde al risultato dell'identità (6). Lasciamo al lettore fornito di software di calcolo più avanzati rendersi conto che qualunque valore assegnerà alla x e alla n risulterà sempre un'identità. In questo caso la parte decrementata è la sola parte reale.

x appartenente all'insieme dei quaternioni (con decremento della sola parte immaginaria i)

```
n = 3
nd = list(range(0,n + 1))
k = 0
```

```

result = 0.00
x = Quaternion(2,4,6,8)
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    kk = Quaternion(2,4 - k,6,8) ** n
    prod = first * comb + kk
    result = result + prod
print(result)

```

x appartenente all'insieme dei quaternioni (con decremento della parte reale e della parte immaginaria i)

```

n = 4
nd = list(range(0, n + 1))
x = Quaternion(5,4,6,9)
k = 0
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    quaternionminusk = Quaternion(5 - k ,4 - k,6,9) ** n
    prod = first * comb * quaternionminusk
    result = result + prod
print(result)

```

Descrizione

Le nuove identità sono esattamente le stesse del caso dei complessi quando vengono decrementate sia la parte reale che la parte immaginaria.

x appartenente all'insieme dei quaternioni (con decremento della parte reale e delle parti immaginarie i e j)

```

getcontext().prec = 1000000000000000000
precision = 1000000000000000000
mp.dps = 1000000000000000000
n = 8
nd = list(range(0, n + 1))
x = Quaternion(5,4,6,9)
k = 0
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    quaternionminusk = Quaternion(5 - k ,4- k,6 - k ,9) ** n

```

```

prod = first * comb * quaternionminusk
result = result + prod
print(result)

```

Quaternion decrementando real, i e j a partire da n = 0 fino a n = 10:

- $\#1.000 + 0.000i + 0.000j + 0.000k$
- $\#1.000 + 1.000i + 1.000j - 0.000k$
- $\# - 2.000 + 4.000i + 4.000j + 0.000k$
- $\# - 30.000 + 6.000i + 6.000j + 0.000k$
- $\# - 168.000 - 96.000i - 96.000j - 0.000k$
- $\#120.000 - 1320.000i - 1320.000j + 0.000k$
- $\#16560.000 - 7200.000i - 7200.000j - 0.000k$
- $\#216720.000 + 65520.000i + 65520.000j + 0.000k$
- $\#685440.000 + 2257920.000i + 2257920.000j + 0.000k$
- $\# - 34473600.000 + 26490240.000i + 26490240.000j - 0.000k$
- $\# - 874540800.000 - 79833600.000i - 79833600.000j + 0.000k$

x appartenente all'insieme dei quaternioni (con decremento della parte reale e delle parti immaginarie i, j e k)

```

n = 4
nd = list(range(0, n + 1))
x = Quaternion(2,4,6,8)
k = 0
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    quaternion_minus_k = Quaternion(2 - k ,4 - k,6 - k ,8 - k) ** n
    prod = first * comb * quaternion_minus_k
    result = result + prod
print(result)
# Quaternion decrementando real, i, j e k a partire da n = 0 fino a n = 6

```

- $\#1.000 + 0.000i + 0.000j + 0.000k$
- $\#1.000 + 1.000i + 1.000j + 1.000k$
- $\# - 4.000 + 4.000i + 4.000j + 4.000k$
- $\#48.000 - 0.000i - 0.000j - 0.000k$
- $\# - 192.000 - 192.000i - 192.000j - 192.000k$
- $\#1920.000 - 1920.000i - 1920.000j - 1920.000k$
- $\#46080.000 - 0.000i - 0.000j - 0.000k$

x appartenente all'insieme dei quaternioni (con decremento delle parti immaginarie i, j e k)

```

n = 2
nd = list(range(0, n + 1))
x = Quaternion(5,4,6,9)
k = 0
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    quaternionminusk = Quaternion(5 ,4 - k,6 -k ,9 - k) ** n
    prod = first * comb * quaternionminusk
    result = result + prod
print(result)

# Quaternion decrementando i, j e k a partire da n = 0 fino a n = 10.

```

- $\#1.000 + 0.000i + 0.000j + 0.000k$
- $\#0.000 + 1.000i + 1.000j + 1.000k$
- $\# - 6.000 - 0.000i - 0.000j - 0.000k$
- $\#0.000 - 18.000i - 18.000j - 18.000k$
- $\#216.000 + 0.000i + 0.000j + 0.000k$
- $\# - 0.000 + 1080.000i + 1080.000j + 1080.000k$
- $\# - 19440.000 - 0.000i - 0.000j - 0.000k$
- $\#0.000 - 136080.000i - 136080.000j - 136080.000k$
- $\#3265920.000 + 0.000i + 0.000j + 0.000k$
- $\#0.000 + 29393280.000i + 29393280.000j + 29393280.000k$
- $\# - 881798400.000 - 0.000i - 0.000j - 0.000k$

x appartenente all'insieme degli ottetti

```

n = 9
nd = list(range(0, n + 1))
x = Octonion(2,3,7,9,5,1,-4,8)
k = 0
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    Octonionminusk = (x - k) ** n
    prod = first * comb * Octonionminusk
    result = result + prod
factorial = math.factorial(n)
difference = result - factorial
print(f'result:result; factorial: factorial; difference:difference')
result:+362880.0000 +0.0000i +0.0000j +0.0000k +0.0000l +0.0000il +0.0000jl
+0.0000kl;
factorial: 362880;
difference:+0.0000 +0.0000i +0.0000j +0.0000k +0.0000l +0.0000il +0.0000jl
+0.0000kl

```

Descrizione

Tale codice esprime l'identità (6) nel caso che la x sia un ottetto e venga decrementata la sola parte reale.

x appartenente all'insieme degli ottetti (con decremento di 7 parti immaginarie)

```

n = 0
nd = list(range(0, n + 1))
x = Octonion(2,3,7,9,5,1, 4,8)
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    octonionminusk = Octonion(2 ,3 -k ,7 - k ,9 - k,5 - k,1 - k, 4 - k,8 -k)
    kk = octonionminusk ** n
    prod = first * comb * kk
    result = result + prod
print(result)

```

Risultati a partire da $n = 0$ fino a $n = 6$

- $\# + 1.0000 + 0.0000i + 0.0000j + 0.0000k + 0.0000l + 0.0000il + 0.0000jl + 0.0000kl$
- $\# + 0.0000 + 1.0000i + 1.0000j + 1.0000k + 1.0000l + 1.0000il + 1.0000jl + 1.0000kl$
- $\# - 14.0000 + 0.0000i + 0.0000j + 0.0000k + 0.0000l + 0.0000il + 0.0000jl + 0.0000kl$
- $\# + 0.0000 - 42.0000i - 42.0000j - 42.0000k - 42.0000l - 42.0000il - 42.0000jl - 42.0000kl$
- $\# + 1176.0000 + 0.0000i + 0.0000j + 0.0000k + 0.0000l + 0.0000il + 0.0000jl + 0.0000kl$
- $\# + 0.0000 + 5880.0000i + 5880.0000j + 5880.0000k + 5880.0000l + 5880.0000il + 5880.0000jl + 5880.0000kl$
- $\# - 246960.0000 + 0.0000i + 0.0000j + 0.0000k + 0.0000l + 0.0000il + 0.0000jl + 0.0000kl$

Descrizione

Quando vengono decrementate le parti immaginarie l'identità si comporta analogamente a quanto visto sopra nel caso dei quaternioni per i , j e k .

x appartenente all'insieme dei sedenioni

```

n = 10
nd = list(range(0, n + 1))
x = Sedenion(2,7,9,8,7,6,5,7,8,4,3,2,9,3,6,5)
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    sedenionminusk = (x - k) ** n
    prod = first * comb * sedenionsminusk
    result = result + prod
factorial = math.factorial(n)

```

```

difference = result - factorial
print(f'result:result; factorial: factorial; difference:difference')
result:(3.6288e+06 0 0 0 0 0 0 0 0 0 0 0 0 0 0);
factorial: 3628800;
difference:(0 0 0 0 0 0 0 0 0 0 0 0 0 0 0)

```

Descrizione

Tale codice esprime l'identità (6) nel caso che la x sia un Sedenione e venga decrementata la sola parte reale. Quando vengono decrementate le parti immaginarie l'identità si comporta analogamente a quanto visto sopra nel caso dei quaternioni per i , j e k .

x appartenente all'insieme dei sedenioni (con decrementando di 6 unità immaginarie)

```

n = 6
nd = list(range(0, n + 1))
x = Sedenion(2,7,9,8,7,6,5,7,8,4,3,2,9,3,6,5)
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    sedenionminusk = Sedenion(2 ,7 ,9 - k,8,7 -k,6 -k , 5 -k ,8 ,4 ,3 -k ,9 ,3,7 - k,6,5, 8)
    ** n
    prod = first * comb * sedenionminusk
    result = result + prod
print(result)

```

Risultati a partire da $n = 0$ fino a $n = 6$.

- $(1, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0)$
- $\#(0, 0, 1, 0, 1, 1, 1, 0, 0, 1, 0, 0, 1, 0, 0, 0)$
- $\#(-12, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0)$
- $(0, 0, -36, 0, -36, -36, -36, 0, 0, -36, 0, 0, -36, 0, 0, 0)$
- $(864, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0)$
- $(0, 0, 4320, 0, 4320, 4320, 4320, 0, 0, 4320, 0, 0, 4320, 0, 0, 0)$
- $\#(-155520, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0)$

x appartenente all'insieme dei sedenioni (con decremento di 14 unità immaginarie)

```

n = 5
nd = list(range(0, n + 1))
x = Sedenion(2,7,9,8,7,6,5,7,8,4,3,2,9,3,6,5)
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    comb = math.comb(n, k)
    sedenionminusk = Sedenion(2 ,7 - k, 9 -k, 8 - k, 7 -k, 6 -k , 5 -k,8 ,4 - k,3 -k,9 -k ,3
-k,7 - k ,6 -k ,5 -k , 8 -k) ** n
    prod = first * comb * sedenionminusk
    result = result + prod
print(result)

```

Risultati a partire da $n = 0$ fino a $n = 5$

- $\#(1, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0)$
- $\#(0, 1, 1, 1, 1, 1, 1, 0, 1, 1, 1, 1, 1, 1, 1, 1)$
- $\#(-28, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0)$
- $\#(0, -84, -84, -84, -84, -84, -84, 0, -84, -84, -84, -84, -84, -84, -84, -84)$
- $(4704, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0)$
- $\#(0, 23520, 23520, 23520, 23520, 23520, 23520, 0, 23520, 23520, 23520, 23520, 23520, 23520, 23520, 23520)$

L'identità con i polinomi di Bernstein

```

def bernstein poly(i, n, t):
    """
    :param i: 0,1,2,3
    :param n: 3
    :param t: 0,1
    :return: comb(n, i) * (t ** (n - i)) * ((1-t) **i)
    """
    return comb(n, i) * (t ** (n - i)) * ((1 - t) ** i)
n = 12
nd = list(range(0,n + 1))
k = 0
x = bernstein poly(2, 3, 0)
result = 0.0
for k in nd:
    first = math.pow(-1, k)

```

```

comb = math.comb(n, k)
second = math.pow((x - k), n)
prod = first * comb * second
result = result + prod
factorial = math.factorial(n)
difference = result - factorial
print(f'result:result; factorial: factorial; difference:difference')
result:479001600.0;
factorial: 479001600;
difference:0.0

```

Descrizione

Da tale codice si evince che l'identità (6) si applica anche al polinomio di Bernstein con parametri scelti arbitrariamente fra quelli possibili.

L'identità con l'insieme di Mandelbrot

```

def mandelbrot(c, max iter=100):
    z = 0
    n = 0
    while abs(z) <= 2 and n < max iter:
        z = z * z + c
        n += 1
    if n == max iter:
        return max iter
    return n + 1 - (math.log2(abs(z)))
c = 0.3 + 0.65j
max iter = 1000
x = mandelbrot(c,max iter)
x = Decimal(x)
n = 3
nd = list(range(0,n + 1))
k = 0
result = 0.0
for k in nd:
    first = math.pow(-1, k)
    first = Decimal(first)
    comb = math.comb(n, k)
    comb = Decimal(comb)
    second = (x - k) ** n
    second = Decimal(second)
    prod = first * comb * second
    prod = Decimal(prod)
    result = Decimal(result)
    result = result + prod
factorial = math.factorial(n)

```

[illegible]

Descrizione

Tale codice determina che l'identità (6) è valida anche per un noto frattale, l'insieme di Mandelbrot. Possiamo supporre che siano molti i frattali dove tale identità (6) è valida. Saremo lieti di sapere se i lettori riusciranno a validare tale identità nei frattali a noi noti.

Conclusioni

La parte formale e la parte computazionale dell'articolo hanno illustrato nuovi risultati che speriamo possano avere utilità pratica e condurre a nuove ricerche teoriche. La **proprietà 3** porta a una generalizzazione del **teorema 1**, dove al posto delle basi delle potenze i -esime dei numeri naturali $k+1$, al variare di k da 0 a n , scriviamo i numeri reali $x-k$. Visto il fatto che l'identità è stata estesa a insiemi sempre più grandi dimensionalmente, si conclude che potrebbe valere in infinite estensioni degli ipercomplessi.

Ringraziamenti

di Giovanni Mare

Ringrazio sentitamente di cuore tutti coloro che hanno contribuito in maniera propulsiva a questo progetto: il mio professore di matematica, il mio professore di statistica, il mio consocio del Mensa, il mio carissimo amico ingegnere, i miei carissimi amici fisici e astrofisici, il mio miglior amico, i miei genitori e la mia famiglia, per avermi sostenuto dal primo all'ultimo giorno. E infine grazie a tutti i miei carissimi amici, e in particolar modo al mio partner, che con pazienza hanno saputo starmi accanto e rendermi felice sempre e comunque. Infinitamente grazie

di Domenico Veca

Desidero esordire ringraziando i miei genitori, la mia ragazza e tutta la mia famiglia: sono stati la mia ancora di speranza nei momenti di incertezza nella continuazione della stesura di questo articolo. Un caro ringraziamento a mia nonna che è stata la persona che mi ha suscitato la curiosità verso la matematica. Uno speciale grazie al mio collaboratore che ha risvegliato in me la voglia di ricercare la bellezza nel mondo dei numeri. Questo articolo è il frutto di una operosa e amichevole cooperazione.

Bibliografia

- [1] "General trattato de' numeri e misure" - Niccolò Tartaglia (1556)
- [2] "The Method of Fluxions" - Sir Isaac Newton (1736)
- [3] "I numeri non sbagliano mai" - Jordan Ellenberg (2015)
- [4] "Pensare in Python" - Allen B. Downey (2016)
- [5] Byjus, <https://byjus.com/maths/pascals-triangle/>
- [6] E.Gupta, Pascal's Triangle–Definition, Properties, Applications and More, 14 May 2024, <https://www.shiksha.com/online-courses/articles/pascals-triangle/>
- [7] Youmath, <https://www.youmath.it>
- [8] Weschool, <https://www.weschool.com/it/>
- [9] Matematicamente.it <https://www.matematicamente.it>
- [10] Wikipedia, https://it.wikipedia.org/wiki/Pagina_principale
- [11] Copilot, <https://copilot.microsoft.com>
- [12] Quora, <https://it.quora.com>

Una procedura per risolvere un sistema lineare di complementarità

A procedure for solving a linear complementarity system

Letizia Pellegrini¹ e Alberto Peretti²

Abstract

We present a procedure for finding a solution of a linear complementarity system. The procedure is based on the theory of the simplex method [3,4] and generates iterative cuts of the relaxed system, until one of the solutions is obtained. This solution can be used in the first step of the method we have described in [5].

Introduzione

Consideriamo un sistema lineare (SL) del tipo

$$\text{SL} \quad \begin{cases} Ax \geq b \\ x \geq 0 \end{cases}$$

dove $A \in \mathbb{R}^{m \times n}$ e $b \in \mathbb{R}^m$.

Come noto, SL può essere visto come la regione ammissibile di un problema di programmazione lineare (PPL). L'algoritmo del simplesso permette di risolvere un PPL e trovare una soluzione ammissibile di SL non presenta difficoltà.

Ora, per preparare la modifica successiva di SL, riformuliamolo nella forma equivalente

$$\text{SL} \quad \begin{cases} Ax + By \geq b \\ x \geq 0, y \geq 0 \end{cases}$$

¹Dipartimento di Scienze Economiche, Università di Verona, email: letizia.pellegrini@univr.it

²Dipartimento di Scienze Economiche, Università di Verona, email: alberto.peretti@univr.it

dove $A, B \in \mathbb{R}^{m \times n}$ e $b \in \mathbb{R}^m$; ovvero abbiamo distinto il vettore variabile in due componenti vettoriali (x, y) con $x = (x_1, x_2, \dots, x_n)$ e $y = (y_1, y_2, \dots, y_n)$.

Ora introduciamo una sostanziale modifica, inserendo un ulteriore vincolo (non lineare) e formuliamo il sistema lineare di complementarità [1,2] (SLC)

$$\text{SLC} \quad \begin{cases} Ax + By \geq b \\ x \geq 0, y \geq 0 \\ \langle x, y \rangle = 0 \end{cases}$$

dove $\langle \cdot, \cdot \rangle$ indica il prodotto interno (scalare) in $\mathbb{R}^n$. Il nuovo vincolo $\langle x, y \rangle = 0$, unitamente ai vincoli di non negatività, pone una cosiddetta condizione di complementarità sulle variabili. Praticamente chiede che, per ogni componente x_i non nulla, debba essere nulla la corrispondente y_i . Si tratta, come detto, di un vincolo non lineare, che quindi fa uscire SLC dalla classe delle regioni ammissibili di un PPL. Vengono a cadere tutte le proprietà di convessità che rendono applicabili le classiche condizioni di ottimalità, sulle quali sostanzialmente si fonda il funzionamento dell'algoritmo del simplesso.

A titolo di semplice esempio geometrico, la Figura 1

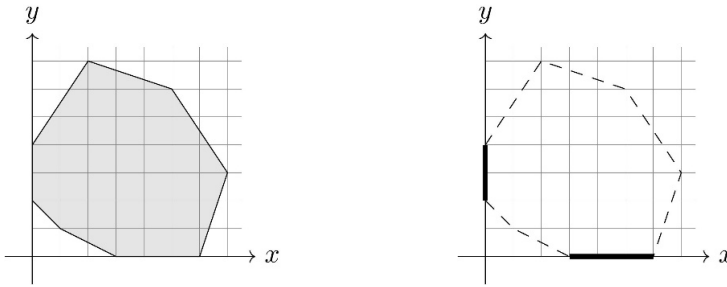


Figura 1

mostra in $\mathbb{R}^2$ a sinistra una tipica regione ammissibile di un SL e a destra, con il tratto molto marcato, le corrispondenti soluzioni del SLC associato.

Vogliamo descrivere una procedura per trovare una soluzione di SLC. Va detto che banalmente una soluzione si può trovare per tentativi enumerando tutte le possibilità di verificare la condizione di complementarità, fissando a turno, sulle variabili, quelle che si devono annullare. Una volta fissate quali componenti si annullano, il problema torna ad essere convesso e quindi vi sono metodi per trovare una soluzione ammissibile. Questo modo, teoricamente possibile, risulta però estremamente pesante quando le dimensioni del problema sono elevate. Si pensi solo che le modalità di annullamento delle n componenti del vettore $(x_1, x_2, \dots, x_n)$ sono tante quante le scritture binarie di lunghezza n , cioè 2^n . Il problema ha quindi una complessità di tipo esponenziale.

Con lo scopo di trovare una soluzione di SLC indichiamo con K l'insieme delle soluzioni di SL e con K_0 l'insieme delle soluzioni di SLC, cioè

$$K = \{(x, y) \in \mathbb{R}^{2n} : Ax + By \geq b, x \geq 0, y \geq 0\}$$

e

$$K_0 = \{(x, y) \in \mathbb{R}^{2n} : Ax + By \geq b, x \geq 0, y \geq 0, \langle x, y \rangle = 0\}.$$

Consideriamo una funzione obiettivo $\downarrow$ *limitata inferiormente* sull'insieme K , come ad esempio $\downarrow(x, y) = \sum_{i=1}^n (x_i + y_i)$. Indichiamo con P il problema

$$\text{P} \quad \begin{cases} \min \ell(x, y) \\ (x, y) \in K_0 \end{cases}$$

e con PR la sua versione rilassata (rimuovendo la condizione di complementarità)

$$\text{PR} \quad \begin{cases} \min \ell(x, y) \\ (x, y) \in K. \end{cases}$$

La forma standard di PR è

$$\text{PR} \quad \begin{cases} \min \ell(x, y) \\ Ax + By + \tilde{I}_m t = b \\ x \geq 0, y \geq 0, t \geq 0 \end{cases}$$

dove $t \in \mathbb{R}^m$ e $b \geq 0$, $\tilde{I}_m$ è una matrice diagonale con $\tilde{I}_m^2 = I_m$ e infine I_m indica la matrice identità in $\mathbb{R}^{m \times m}$. $\tilde{I}_m$ è pertanto una matrice diagonale che ha ± 1 sulla diagonale principale. Infatti, per avere $b \geq 0$, come prevede la forma standard, è necessario che le variabili scarto t_i figurino nel sistema come $\pm t_i$.

Supponiamo infine che tutti i vertici della regione ammissibile di PR corrispondano a soluzioni di base non degeneri. Nella sezione seguente descriviamo una procedura per trovare una soluzione ammissibile di SLC. Per illustrare la procedura forniamo poi due esempi numerici.

Descrizione della procedura

Risolviamo PR e sia $(\bar{x}, \bar{y}, \bar{t})$ una sua soluzione ottima di base non degenera. Se $\langle x, y \rangle = 0$, allora sia $(\bar{x}, \bar{y})$ è una soluzione di SLC e la procedura termina. Se invece $\langle x, y \rangle > 0$, sia $(\bar{z}_D, \bar{z}_N)$ la scomposizione di $\bar{z} = (\bar{x}, \bar{y}, \bar{t})$ nelle componenti rispettivamente di base e non di base; sia poi (D, N) la corrispondente scomposizione in matrice di base e non di base di $(A, B, \tilde{I}_m)$ (supponiamo, senza perdere generalità, che le componenti di base siano le prime m). Abbiamo allora

$$\bar{z} = (\bar{z}_D = D^{-1}b - D^{-1}N\bar{z}_N, \bar{z}_N = 0).$$

Indichiamo con $(\beta_{i0}), i = 1, \dots, m$, il vettore (di m componenti) $D^{-1}b$, e con $(\beta_{ij}), i = 1, \dots, m, j = 1, \dots, 2n$, la matrice $(m \times n)$ $D^{-1}N$. Per ogni $k \in I_N = \{m+1, \dots, m+2n\}$ sia z_k^{sup} definito, come nell'algoritmo del simplesso [3,4], da

$$z_k^{\sup} = \begin{cases} +\infty, & \text{se } \beta_{ik} \leq 0, i = 1, \dots, m \\ \min_{\beta_{ik}} \left\{ \frac{\beta_{i0}}{\beta_{ik}} \right\}, & \text{altrimenti} \end{cases} \quad (1)$$

Ricordiamo che $z_k^{\sup}$ è il valore massimo che può assumere la k -esima componente non di base di z_N e che mantiene l'ammissibilità della soluzione. Nell'ipotesi che l'insieme ammissibile di PR sia limitato, abbiamo che $z_k^{\sup} < +\infty, \forall k \in I_N$; inoltre, se $\bar{z}$ è una soluzione di base non degenera, allora $z_k^{\sup} > 0, \forall k \in I_N$.

Definiamo ora il seguente vettore $\hat{z} \in \mathbb{R}^{m+2n}$:

$$\begin{cases} \hat{z}_j = \bar{z}_j, j = 1, \dots, m \\ (\hat{z}_{m+1}, \dots, \hat{z}_{m+2n}) = (\lambda_1 z_{m+1}^{\sup}, \dots, \lambda_{2n} z_{m+2n}^{\sup}) = \sum_{s=1}^{2n} \lambda_s z^{m+s} \end{cases} \quad (2)$$

con $\lambda_s \geq 0, \sum_{s=1}^{2n} \lambda_s < 1$, e

$$z^{m+1} = (z_{m+1}^{\sup}, 0, \dots, 0), \dots, z^{m+2n} = (0, 0, \dots, z_{m+2n}^{\sup}).$$

Sia ora H_0 l'iperpiano (unico) passante per i $2n$ punti $z^{m+s}, s = 1, \dots, 2n$. Si ha

$$H_0 = \left\{ (z_{m+1}, z_{m+2}, \dots, z_{m+2n}) \in \mathbb{R}^{2n} : \frac{z_{m+1}}{z_{m+1}^{\sup}} + \frac{z_{m+2}}{z_{m+2}^{\sup}} + \dots + \frac{z_{m+2n}}{z_{m+2n}^{\sup}} = 1 \right\}.$$

Proposizione

Se $(z_{m+1}, z_{m+2}, \dots, z_{m+2n})$ appartiene all'insieme

$$H_{\geq 0}^- = \left\{ (z_{m+1}, z_{m+2}, \dots, z_{m+2n}) \in \mathbb{R}^{2n} : \frac{z_{m+1}}{z_{m+1}^{\sup}} + \frac{z_{m+2}}{z_{m+2}^{\sup}} + \dots + \frac{z_{m+2n}}{z_{m+2n}^{\sup}} < 1 \right\}$$

e (x, y, t) è la corrispondente soluzione ammissibile di PR, allora $\langle x, y \rangle > 0$.

Dimostrazione. Se $(z_{m+1}, \dots, z_{m+2n}) \in H_{\geq 0}^-$, allora, ponendo $\frac{z_{m+s}}{z_{m+s}^{\sup}} = \lambda_s, s = 1, \dots, 2n$, si ha

$$(z_{m+1}, \dots, z_{m+2n}) = (\lambda_1 z_{m+1}^{\sup}, \dots, \lambda_{2n} z_{m+2n}^{\sup}) = \sum_{s=1}^{2n} \lambda_s z^{m+s},$$

$$\lambda_s \geq 0, s = 1, \dots, 2n, \text{ e } \sum_{s=1}^{2n} \lambda_s = \sum_{s=1}^{2n} \lambda_s \frac{z_{m+s}}{z_{m+s}^{\sup}} < 1.$$

La tesi segue dalla (2) e dal fatto che i vettori $\hat{z} = (\hat{x}, \hat{y}, \hat{t})$ definiti in (2) soddisfano la condizione $\langle \hat{x}, \hat{y} \rangle > 0$. Infatti, basta osservare che nella soluzione $\bar{z} = (\bar{x}, \bar{y}, \bar{t})$ le componenti nulle ora, nel vettore $\hat{z} = (\hat{x}, \hat{y}, \hat{t})$ sono positive, mentre tutte le altre componenti rimangono inalterate in $\hat{z}$.

□

Pertanto, dalla Proposizione 1, se aggiungiamo la disuguaglianza

$$\frac{z_{m+1}}{z_{m+1}^{\sup}} + \frac{z_{m+2}}{z_{m+2}^{\sup}} + \dots + \frac{z_{m+2n}}{z_{m+2n}^{\sup}} \geq 1 \quad (3)$$

come nuovo vincolo nella regione ammissibile di PR, questa è un taglio nell'insieme ammissibile K che non esclude alcuna soluzione di P che soddisfa la condizione di complementarità $\langle x, y \rangle = 0$. Aggiungiamo la disuguaglianza (3) alla regione ammissibile di PR e sia PR_1 il problema così ottenuto. Riappliciamo poi la procedura al nuovo problema PR_1 .

Osservazione 1

Abbiamo ottenuto il precedente risultato nell'ipotesi che la regione ammissibile di PR sia limitata. Se non fosse così, in (1) almeno uno dei valori $z_k^{\sup}$ potrebbe essere $+\infty$. Supponiamo che accada per i primi p , con $m+1 \leq p < m+2n$. Allora sostituiamo il taglio (3) con

$$\frac{z_{m+p+1}}{z_{m+p+1}^{\sup}} + \dots + \frac{z_{m+2n}}{z_{m+2n}^{\sup}} \geq 1. \quad (4)$$

Se $p = m+2n$, cioè tutti i valori $z_k^{\sup}$ sono $+\infty$, significa che siamo in una soluzione ottima di PR e non possiamo effettuare nessun taglio. Pertanto la procedura termina: se la soluzione non soddisfa la complementarità significa che non c'è una soluzione ammissibile del SLC.

L'altra ipotesi fatta su PR è che tutte le sue soluzioni di base siano non degeneri. Se $(\bar{x}, \bar{y}, \bar{t})$ è una soluzione di base degenera di PR dobbiamo applicare uno dei metodi classici per gestire la degenerazione; per esempio, usare un criterio lessicografico di entrata in base o una ϵ -perturbation. In questo modo in (1) otteniamo $z_k^{\sup} > 0, \forall k \in I_N$.

Osservazione 2

Ad ogni passo della procedura i tagli (3) oppure (4) escludono un sottoinsieme di K che non contiene soluzioni (x, y) tali che $\langle x, y \rangle = 0$. La procedura termina quando risolvendo uno dei PR troviamo una soluzione ottima $(\bar{x}, \bar{y}, \bar{t})$ tale che $\langle \bar{x}, \bar{y} \rangle = 0$ oppure quando non c'è una soluzione ammissibile, come indicato nell'Osservazione 1.

Esempi

In questo paragrafo proponiamo un paio di esempi per illustrare la procedura descritta in precedenza. Gli esempi hanno una dimensione $n = 1$ perché così possiamo avere una rappresentazione geometrica in $\mathbb{R}^2$.

Esempio 1. Consideriamo il seguente SLC con $n = 1, m = 5$:

$$\begin{cases} \min(x + y) \\ x \geq 1 \\ x + 2y \geq 5 \\ x + 4y \geq 7 \\ -x + y \leq 3 \\ x \geq 0, y \geq 0 \\ \langle x, y \rangle = 0 \end{cases}$$

e il relativo problema P con $\uparrow(x, y) = x + y$. Se introduciamo il problema rilassato, ottenuto togliendo la condizione di complementarità, la sua regione ammissibile è rappresentata in Figura 2, dove è indicato in grassetto tratteggiato l'insieme di livello della funzione obiettivo corrispondente al suo valore minimo.

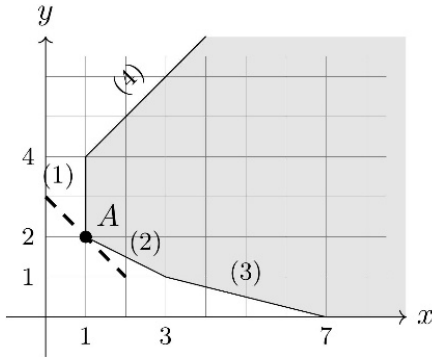


Figura 2

$$\begin{aligned} & \min(x + y) \\ (1) \quad & x \geq 1 \\ (2) \quad & x + 2y \geq 5 \\ (3) \quad & x + 4y \geq 7 \\ (4) \quad & -x + y \leq 3 \\ & x \geq 0, y \geq 0. \end{aligned}$$

La forma standard è

$$\begin{cases} \min(x + y) \\ x - t_1 = 1 \\ x + 2y - t_2 = 5 \\ x + 4y - t_3 = 7 \\ -x + y + t_4 = 3 \\ x \geq 0, y \geq 0, t_1, t_2, t_3, t_4 \geq 0. \end{cases}$$

La prima e ultima tabella delle iterazioni del simplesso sono [3,4]

x	y	t_1	t_2	t_3	t_4		x	y	t_1	t_2	t_3	t_4	
1	0	-1	0	0	0	1	1	0	-1	0	0	0	1
1	2	0	-1	0	0	5	0	1	$1/2$	$-1/2$	0	0	2
1	4	0	0	-1	0	7	0	0	1	-2	1	0	2
-1	1	0	0	0	1	3	0	0	$-3/2$	$1/2$	0	1	2
							0	0	$1/2$	$1/2$	0	0	

Dalla tabella a destra abbiamo che la soluzione ottima è

$$(\bar{x}, \bar{y}, \bar{t}_1, \bar{t}_2, \bar{t}_3, \bar{t}_4) = (1, 2, 0, 0, 2, 2)$$

in cui la partizione in componenti di base e non di base è

$$\bar{z}_D = (\bar{x}, \bar{y}, \bar{t}_3, \bar{t}_4) \quad \text{e} \quad \bar{z}_N = (\bar{t}_1, \bar{t}_2).$$

La soluzione corrisponde al vertice $A = (1, 2)$ in Figura 2. Chiaramente la condizione di complementarità non è soddisfatta.

Dalla stessa tabella possiamo ricavare l'informazione relativa al primo taglio.

Le due variabili che possono entrare in base sono t_1 e t_2 . I loro valori massimi, che mantengono l'appartenenza alla regione ammissibile, applicando la (1), sono

$$z_{t_1}^{\sup} = \min(4, 2) = 2 \quad \text{e} \quad z_{t_2}^{\sup} = \min(4) = 4.$$

Questi due valori definiscono la disequazione (vedi (3))

$$\frac{t_1}{2} + \frac{t_2}{4} \geq 1 \quad \text{o equivalentemente} \quad 2t_1 + t_2 \geq 4.$$

Possiamo facilmente ricavare la forma equivalente in termini delle variabili originarie, usando le equazioni della forma standard $t_1 = x - 1$ e $t_2 = x + 2y - 5$, e otteniamo $3x + 2y \geq 11$.

Se aggiungiamo la disuguaglianza $3x + 2y \geq 11$ all'insieme ammissibile, i vincoli (1) e (2) in Figura 2, cioè $x \geq 1$ e $x + 2y \geq 5$ diventano ridondanti; abbiamo quindi il nuovo problema indicato a fianco della Figura 3.

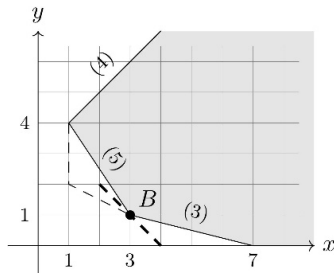


Figura 3

$$\begin{aligned} \min(x+y) \\ (3) \quad & x + 4y \geq 7 \\ (4) \quad & -x + y \leq 3 \\ (5) \quad & 3x + 2y \geq 11 \\ & x \geq 0, y \geq 0. \end{aligned}$$

Dalla Figura 3 possiamo vedere che la soluzione ottima è ora $B = (3, 1)$. La forma standard è la seguente

$$\begin{cases} \min(x + y) \\ x + 4y - t_3 = 7 \\ -x + y + t_4 = 3 \\ 3x + 2y - t_5 = 11 \\ x \geq 0, y \geq 0, t_3, t_4, t_5 \geq 0. \end{cases}$$

La prima e ultima tabella delle iterazioni del simplesso sono

x	y	t_3	t_4	t_5		x	y	t_3	t_4	t_5	
1	4	-1	0	0	7	1	0	$2/_{10}$	0	$-4/_{10}$	3
-1	1	0	1	0	3	0	1	$-3/_{10}$	0	$1/_{10}$	1
3	2	0	0	-1	11	0	0	$5/_{10}$	1	$-5/_{10}$	5
						0	0	$1/_{10}$	0	$11/_{10}$	

La soluzione ottima è quindi

$$(\bar{x}, \bar{y}, \bar{t}_3, \bar{t}_4, \bar{t}_5) = (3, 1, 0, 5, 0)$$

in cui la partizione in componenti di base e non di base è

$$\bar{z}_D = (\bar{x}, \bar{y}, \bar{t}_4) \quad \text{e} \quad \bar{z}_N = (\bar{t}_3, \bar{t}_5).$$

La soluzione corrisponde al vertice $B = (3, 1)$ in Figura 3. La condizione di complementarità non è soddisfatta.

Dalla stessa tabella possiamo ricavare l'informazione relativa al secondo taglio. Le due variabili che possono entrare in base sono t_3 e t_5 . I loro valori massimi, che mantengono l'appartenenza alla regione ammissibile, applicando la (1), sono

$$z_{t_3}^{\sup} = \min(15, 10) = 10 \quad \text{e} \quad z_{t_5}^{\sup} = \min(10) = 10.$$

Questi due valori definiscono la disequazione (vedi (3))

$$\frac{t_3}{10} + \frac{t_5}{10} \geq 1 \quad \text{o equivalentemente} \quad t_3 + t_5 \geq 10.$$

Possiamo facilmente ricavare la forma equivalente in termini delle variabili originali, usando le equazioni della forma standard $t_3 = x + 4y - 7$ e $t_5 = 3x + 2y - 11$, e otteniamo $2x + 3y \geq 14$.

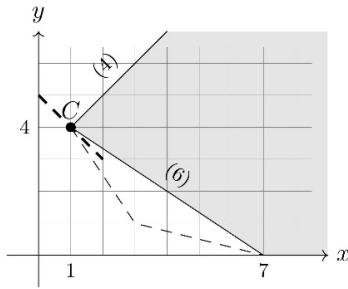


Figura 4

$$\begin{aligned} & \min(x + y) \\ (4) \quad & -x + y \leq 3 \\ (6) \quad & 2x + 3y \geq 14 \\ & x \geq 0, y \geq 0. \end{aligned}$$

La forma standard è

$$\begin{cases} \min(x + y) \\ -x + y + t_4 = 3 \\ 2x + 3y - t_6 = 14 \\ x \geq 0, y \geq 0, t_4, t_6 \geq 0. \end{cases}$$

La prima e ultima tabella delle iterazioni del semplice sono

x	y	t_4	t_6	
-1	1	1	0	3
2	3	0	-1	14

x	y	t_2	t_7	
1	0	$-6/10$	$-2/10$	1
0	1	$4/10$	$-2/10$	4
0	0	$2/10$	$4/10$	

La soluzione ottima è quindi

$$(\bar{x}, \bar{y}, \bar{t}_4, \bar{t}_6) = (1, 4, 0, 0)$$

in cui la partizione in componenti di base e non di base è

$$\bar{z}_D = (\bar{x}, \bar{y}) \quad \text{e} \quad \bar{z}_N = (\bar{t}_4, \bar{t}_6).$$

La soluzione corrisponde al vertice $C = (1, 4)$ in Figura 4. La condizione di complementarità non è soddisfatta.

Dalla stessa tabella possiamo ricavare l'informazione relativa al terzo taglio.

Le due variabili che possono entrare in base sono t_4 e t_6 . I loro valori massimi, che mantengono l'appartenenza alla regione ammissibile, applicando la (1), sono

$$z_{t_4}^{\sup} = \min(10) = 10 \quad \text{e} \quad z_{t_6}^{\sup} = +\infty,$$

che definiscono la disequazione (vedi (4))

$$\frac{t_4}{10} \geq 1 \quad \text{o equivalentemente} \quad t_4 \geq 10.$$

Ricaviamo la forma equivalente in termini delle variabili originarie, usando l'equazione $t_4 = 3 + x - y$, e otteniamo $x - y \geq 7$.

Se aggiungiamo la disuguaglianza $x - y \geq 7$ all'insieme ammissibile, i vincoli (4) e (6) in Figura 4 diventano ridondanti e abbiamo il nuovo problema indicato a fianco della Figura 5.

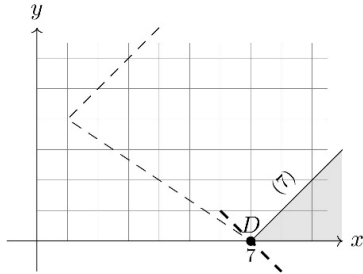


Figura 5

$$(7) \quad \begin{cases} \min(x + y) \\ x - y \geq 7 \\ x \geq 0, y \geq 0. \end{cases}$$

La forma standard è

$$\begin{cases} \min(x + y) \\ x - y - t_7 = 7 \\ x \geq 0, y \geq 0, t_7 \geq 0. \end{cases}$$

La relativa tabella del simplesso è

x	y	t_7	
1	-1	-1	7
1	-1	-1	

che banalmente fornisce la soluzione ottima

$$(\bar{x}, \bar{y}, \bar{t}_7) = (7, 0, 0),$$

corrispondente al vertice $D = (7, 0)$ in Figura 5, per la quale è soddisfatta la condizione di complementarità.

Esempio 2. In questo secondo esempio vogliamo mostrare una situazione in cui l'eliminazione dei vincoli (quelli che prima erano ridondanti) può portare alla non corretta individuazione di una soluzione ammissibile. Consideriamo il seguente esempio di problema P, con $n = 1, m = 2$:

$$\begin{cases} \min(x+y) \\ x+2y \geq 3 \\ x+3y \leq 4 \\ x \geq 0, y \geq 0 \\ \langle x, y \rangle = 0. \end{cases}$$

La regione ammissibile del problema rilassato è rappresentata in Figura 6 e a fianco è scritto il problema rilassato.

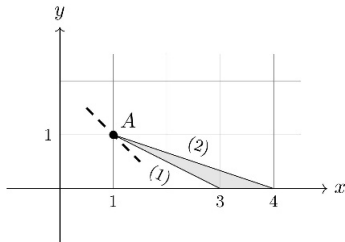


Figura 6

$$\begin{aligned} & \min(x+y) \\ (1) \quad & x+2y \geq 3 \\ (2) \quad & x+3y \leq 4 \\ & x \geq 0, y \geq 0. \end{aligned}$$

La forma standard è

$$\begin{cases} \min(x+y) \\ x+2y-t_1=3 \\ x+3y+t_2=4 \\ x \geq 0, y \geq 0, t_1, t_2 \geq 0. \end{cases}$$

La prima e ultima tabella delle iterazioni del simpleso sono

x	y	t_1	t_2	
1	2	-1	0	3
1	3	0	1	4

x	y	t_1	t_2	
1	0	-3	-2	1
0	1	1	1	1
0	0	2	1	

La soluzione ottima è quindi

$$(\bar{x}, \bar{y}, \bar{t}_1, \bar{t}_2) = (1, 1, 0, 0)$$

e la partizione in componenti di base e non di base è

$$\bar{z}_D = (\bar{x}, \bar{y}) \quad \text{e} \quad \bar{z}_N = (\bar{t}_1, \bar{t}_2).$$

La soluzione corrisponde al vertice $A = (1, 1)$ in Figura 6. La condizione di complementarità non è verificata. L'applicazione di (1) ci permette di determinare

$$z_{t_1}^{\sup} = \min(1) = 1 \quad \text{e} \quad z_{t_2}^{\sup} = \min(1) = 1$$

che definiscono la disequazione $t_1 + t_2 \geq 1$, che in forma equivalente in termini delle variabili originarie (dalla forma standard $t_1 = x + 2y - 3$ e $t_2 = 4 - x - 3y$) diventa $y \leq 0$. Il nuovo problema è

$$\begin{cases} \min(x + y) \\ x + 2y - t_1 = 3 \\ x + 3y + t_2 = 4 \\ y + t_3 = 0 \\ x \geq 0, y \geq 0, t_1, t_2, t_3 \geq 0. \end{cases}$$

La soluzione ottima è il punto $(\bar{x}, \bar{y}, \bar{t}_1, \bar{t}_2, \bar{t}_3) = (3, 0, 0, 1, 0)$, che è soluzione di base degenera.

Si osservi che in questo esempio, a differenza di quanto accade nell'Esempio 1, il taglio non produce la ridondanza di alcuni dei vincoli e quindi il successivo problema ha dimensione maggiore del precedente.

Potrebbe essere interessante indagare a questo punto sul possibile riconoscimento automatico di queste situazioni, in vista della scrittura di un algoritmo generale per la soluzione di un sistema di complementarità.

Bibliografia

- [1] Cottle R.W.: Linear Complementarity Problem. In: Floudas C., Pardalos P. (eds) Encyclopedia of Optimization, Springer, Boston, MA (2008).
- [2] Cottle R.W., Pang J.S., Stone R.E.: The linear complementarity problem, Academic Press, San Diego, CA, 1992; reprint, SIAM Classics in Applied Mathematics, Vol.60, Philadelphia (2009).
- [3] Dantzig G.B.: Linear Programming and Extensions, Princeton University Press, Princeton N.Y. (1963).
- [4] Fischetti M.: Lezioni di Ricerca Operativa, Ed. Libreria Progetto Padova, Padova (2018).
- [5] Mastroeni G., Pellegrini L., Peretti A.: On linear problems with complementarity constraints, Optimization Letters, Vol. 16: pp. 2241–2260 (2022).

Proprietà elastiche e problemi di misura ottimale: un approccio didattico

Elastic properties and optimal measurement problems: a Didactic Approach

Corrado Binetti¹

Abstract

Partendo dal punto di vista di de Finetti tipicamente archimedeo, intriso di quel "fusionismo" di cui tanto caldeggiò l'utilizzo nell'insegnamento scolastico, si è voluto trattare l'argomento in modo euristico, allo scopo di mettere in risalto come, da un punto di vista didattico, sia opportuno non tanto propinare agli alunni direttamente formule o leggi più o meno complicate, quanto fornire agli stessi dei mezzi di osservazione della realtà, quali i modelli, che possano condurli ad intravedere o scoprire le leggi medesime. Lo scopo delle esperienze, rivolte a ragazzi di una quinta classe di una scuola superiore, è quello di illustrare i concetti di curva geodetica e superficie minimale e i problemi isoperimetrici e di Plateau, attraverso le proprietà elastiche della gomma e delle soluzioni saponose.

Problema di Didone: *Didone era la figlia di un re di Tiro che secondo la leggenda era sposata al proprio zio Acerba. Questo fu ucciso a causa delle sue molte ricchezze e così Didone si trovò costretta a fuggire. Approdò a Cipro e da lì salpò verso le coste dell'Africa di fronte alla Sicilia. Approdata sulla costa, chiese al signore del luogo di poter acquistare un po' di terra lungo la spiaggia: un pezzo non più grande di quanto potesse essere cinto con una pelle di bue. Egli acconsentì a questa modesta richiesta da parte di una bella signora e le offrì generosamente una grande pelle. Didone si fece furba due volte, e riuscì ad ottenere ben più terra di quanta il signore del luogo avesse immaginato. Tagliò infatti la pelle in sottilissime striscioline che legò insieme in modo da formare una fune. Si trovò poi di fronte al problema: Qual 'e la figura di area massima che si può circondare con una fune di data lunghezza i cui estremi poggiano su una linea retta? (Si suppone che la spiaggia sia rettilinea) Didone trovò la soluzione del problema e divenne, in parte a causa della sua soluzione ben riuscita, la fondatrice e regina della prospera città di Cartagine.*

(http://web.math.unifi.it/pls/laboratori/minimo_massimo/soluzioni3.pdf)

¹ Istituto Alberghiero di Molfetta (BA) - corradoabinetti@libero.it, corradoabinetti76@gmail.com

Premessa

Le esperienze proposte agli studenti, che completano l'argomento "***I concetti di forma e misura nella geometria moderna***", coinvolgono discipline di grosso spessore culturale: matematica, fisica e chimica e prendono forma in *laboratori di ricerca azione* utilizzando materiale povero e di facile reperibilità. La fase esperienziale è preceduta da una fase teorica, in modalità multidisciplinare, allo scopo di richiamare e/o introdurre nozioni di matematica, fisica e chimica, veicoli indispensabili per le esperienze stesse.

Teoria Matematica

D) E' noto che la lunghezza di una curva C nello spazio euclideo è il limite delle lunghezze di linee spezzate Ln che "approssimano" la curva.

E' semplice immaginare (e dimostrare) che nello spazio euclideo, il "percorso più breve" fra due punti P e Q e' determinato dal segmento $\overline{PQ}$ ma il problema diventa notevolmente più complicato se il percorso e' soggetto a "vincoli".

Per determinare il "percorso più breve" fra due punti P e Q su una superficie S nello spazio, si può cercare, fra tutte le curve contenute in S che hanno P e Q per estremi, la curva C di lunghezza minima.

Per una proficua indagine matematica di questo problema è opportuno analizzare il problema da un punto di vista leggermente più generale: studiare le curve geodetiche su una superficie.

Definizione. Una curva C contenuta in una superficie S si dice **geodetica** se per ogni punto P della curva e per ogni curva C' , contenuta in S e che coincide con la curva C tranne che in prossimità del punto P , si ha che la lunghezza di C' è maggiore della lunghezza di C .

(figura 1)



Da ciò si deduce che la proprietà di essere una geodetica è una proprietà "locale" della curva. Tuttavia, una volta note le geodetiche di una superficie, è generalmente molto semplice risolvere il problema del percorso più breve, perché le soluzioni vanno cercate fra le curve geodetiche che passano per i punti assegnati: in molti casi queste ultime sono in numero finito. Ad esempio, le geodetiche su una sfera sono gli archi di cerchio massimo, cioè le curve ottenute come intersezione della sfera con un piano passante per il centro. Dati due punti sulla sfera, perciò, il percorso più breve è determinato da uno dei due archi di cerchio massimo che hanno per estremi i punti (nel solo caso di punti antipodali ci sono infiniti percorsi "più brevi possibile").

II) E' noto che l'area di una superficie S nello spazio euclideo è il limite di aree di poliedri P_n che "approssimano" la superficie.

Per le superfici, si possono studiare problemi simili a quelli visti per le curve. Analogamente al concetto di geodetica è quello di *superficie minimale*.

Definizione. Una superficie S nello spazio si dice **minimale** se per ogni punto P della curva e per ogni superficie S' che coincide con la superficie S tranne che in prossimità del punto P , si ha che l'area di S' è maggiore dell'area di S .

(figura 2)



Analogamente al problema del "percorso più breve" per le curve è il **Problema di Plateau**: Data una curva chiusa C nello spazio, qual è la superficie S di area minima che ha tale curva come bordo? Esso è spesso chiamato il "problema delle bolle di sapone", perché le **bolle di sapone** forniscono una "evidenza sperimentale" dell'esistenza di tali superfici.

Problema isoperimetrico nel piano: Qual è tra tutte le figure piane di assegnato perimetro, quella di area massima? Si dimostra che il cerchio risolve il problema.

Problema isoperimetrico nello spazio: Qual è la figura tridimensionale che, a parità di volume, ha la minore superficie esterna? Si dimostra che la sfera risolve il problema.

Teoria: Fisica - Chimica

I) Un elastico è un pezzo di gomma di forma pressoché lineare. Le proprietà della gomma fanno sì che, entro certi limiti, esso risponda alle elongazioni con una forza di richiamo omogenea (se il materiale è omogeneo) nella direzione della sua estensione: tende a "ritirarsi su se stesso".

II) Una soluzione saponosa (nel caso in esame: 10 parti acqua distillata, 2 parti detersivo per stoviglie, 1 parte glicerina) si può ottenere in diversi modi.

La proprietà che interessa osservare è che in una tale soluzione, sono disciolte in acqua molecole con due componenti. Una componente "idrofila", che tende a legarsi alle molecole di acqua, ed una componente "idrofoba", che tende ad isolarsi dall'acqua.

In presenza di un oggetto estraneo (ad esempio del grasso) in una soluzione saponosa, la componente idrofoba si lega all'oggetto, mentre quella idrofila si lega all'acqua (questo è il motivo per il quale l'acqua saponata elimina il grasso dai piatti durante il lavaggio!).

In generale, una soluzione di acqua e sapone ha una bassa tensione superficiale. Si può aumentare tale tensione aggiungendo dei tensioattivi (come la glicerina o lo zucchero). La glicerina ha anche la proprietà di rendere più denso il fluido, e quindi l'evaporazione della soluzione è rallentata.

Laboratorio di ricerca azione (I)

Materiali: [pallone da basket, elastici, pennarelli]

- Si disegnano due punti sul pallone da basket con il pennarello, e si chiede allo studente di disegnare quello che lui reputa il percorso più breve, sulla superficie del pallone, per andare da un punto all'altro.
- Si verifica poi la risposta con un elastico. Si fissano le estremità nei due punti (reggendoli con le dita), si fa tendere l'elastico e poi lo si lascia andare.
- In virtù delle proprietà elastiche su descritte, l'elastico si disporrà lungo una geodetica della sfera, e dunque lungo un arco di cerchio massimo.



Osservazioni:

Guardando una sfera, essa ci appare come un disco. Le geodetiche ci appaiono invece come delle ellissi con asse maggiore sul bordo di tale disco. Questa distorsione prospettica rende a volte difficile immaginare quale sia "il percorso più breve" fra due punti, per chi non lo conosca.

Pertanto, è opportuno scegliere gli estremi dell'arco in posizioni che in prospettiva ci appaiono vicino al bordo ma non su un diametro del disco che vediamo.

Con lo stesso meccanismo si possono determinare le geodetiche di numerose altre superfici, come per esempio gli ellissoidi o gli ovoidi... a patto di averne dei modelli piuttosto solidi a disposizione. La condizione fondamentale è che queste abbiano "curvatura positiva" (siano convesse). Nel caso di superfici a curvatura negativa (come per esempio un iperboloide iperbolico), l'esperienza non funziona perchè l'elastico si tende lungo una retta e abbandona la superficie.

Laboratorio di ricerca azione (II)

Materiali: [soluzione saponosa, filo di ferro, carta adesiva a nastro, spago, bacinella, pinza e vasetti di vetro].

Usando filo di ferro, pinza e vasetti di vetro, si costruiscono facilmente supporti con la forma di una circonferenza. Altri supporti si creano curvando il filo di ferro. Per assicurare un maggior "grip" della soluzione saponosa sui supporti costruiti con il filo di ferro, è opportuno avvolgere della carta adesiva lungo il filo. Per assicurare che una maggior quantità di soluzione resti legata al filo-supporto, è infine opportuno avvolgere dello spago intorno al filo come un'elica.



Con questi accorgimenti si possono realizzare le seguenti esperienze:

A) Si soffiano alcune bolle di sapone e si nota che queste tendono ad assumere una forma sferica. Difatti la superficie della bolla tende a ritirarsi su se stessa, ma non può schiacciare completamente l'aria contenuta nella bolla. Il fatto che la forma si stabilizzi su una sfera è una prova pratica del principio isoperimetrico.



Le bolle molto grandi hanno in realtà un rigonfiamento sul fondo. Questo è dovuto al peso dell'acqua che scivola lungo la bolla e si deposita nella parte inferiore.

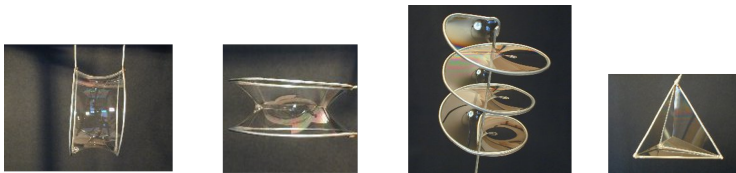


(*)

B) Si costruiscono con il filo di ferro curve di forma diversa, immergendole poi nella soluzione saponosa. Estraeando il modello dalla soluzione, si nota che si è formata una pellicola

Dal punto di vista fisico, la formazione della pellicola fluida sul contorno di filo di ferro si deve sia alla forma del filo, sia alla tensione superficiale del film che tende a ridurne il più possibile l'estensione.

Dal punto di vista matematico, la forma della pellicola deve essere una superficie minimale. La lamina di sapone forma una superficie la cui area è la minima possibile tra quelle aventi quel dato contorno (problema di Plateau); si ottengono, così delle forme che per il principio di minima energia che la natura sceglie sono le migliori possibili.



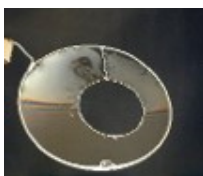
Per quanto riguarda il **principio isoperimetrico** nel piano, si può fornire un primo “incipit” sperimentale di questa proprietà: dando una forma arbitraria ad un nastro flessibile e riempiendolo di palline, si osserva che esso tende ad assumere una forma circolare.



E continuare con **le bolle di sapone**:

Si forma con un supporto una lamina saponata piana, si posiziona delicatamente un “cappio” molto sottile (ad es. un filo di cotone) sulla lamina, si scoppia la parte della lamina superficie all’interno del cappio, si ottiene un buco di forma perfettamente circolare.

La lamina saponata ha massimizzato l’area del buco .



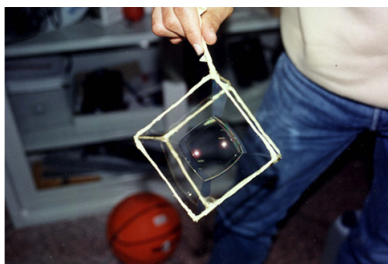
C) Si costruisce con il filo di ferro una struttura, formata dagli spigoli di un cubo. Se si immerge tale struttura nella soluzione saponosa e la si estrae con cautela, si può creare una bolla di forma pressoché cubica legata agli spigoli del cubo da alcuni pezzi di superficie piana.



Osservazioni:

La bolla cubica non contraddice il principio isoperimetrico, perché è soggetta a vincoli esterni, determinati dall'attaccamento agli spigoli del cubo.

Si tratta di una condizione ibrida fra i casi A) e B) analizzati in precedenza.

**Conclusioni**

Le esperienze mostrate sono principalmente illustrative e l'approccio è prevalentemente intuitivo. Tuttavia si riesce a coinvolgere gli studenti in una indagine:

- Utile [perché di evidente interesse in problemi pratici: costruire una rete telefonica che sia la più breve possibile, una rete stradale più breve che collega un certo numero di città, una superficie di area minima che ha come contorno una fissata curva dello spazio (tende di Otto Frei: padiglione dell'Expo 67 a Montreal, stadio di Monaco)]
- Interessante (perché semplice da capire, ma non banale da risolvere)
- Dall'indubbio fascino - perchè no- artistico.

Vorrei concludere con una frase di **Mark Twain**: "Una bolla di sapone è la cosa più bella, e la più elegante che ci sia in natura... Mi chiedo quanto sarebbe necessario per comprare una bolla di sapone se al mondo ne esistesse soltanto una".

(*) Le foto ritraggono un mio carissimo amico fraterno, nonché collega, il Prof. Giulio Minervini, già ricercatore di Geometria presso l'Università A. Moro di Bari, prematuramente scomparso, a cui questo articolo è dedicato.

Matematica quantistica

Quantum mathematics

Luca Granieri

Abstract

Si fornisce un compendio introduttivo sulla matematica della meccanica quantistica a partire dalla formulazione Lagrangiana-Hamiltoniana della meccanica classica.

La rivoluzione quantistica novecentesca costituisce una grande conquista del pensiero scientifico sul terreno di un ripensamento fondamentale dell'indagine scientifica della realtà. Se ci venisse chiesto di identificare i caratteri più rivoluzionari in rottura con la fisica classica in genere si penserebbe all'indeterminazione e alla probabilità coinvolta strutturalmente nei meandri più intimi della materia. Certamente, ma incertezza e probabilità sono contemplate anche nella meccanica classica. Che cosa allora distingue nettamente la meccanica quantistica da quella classica? E perché la formulazione matematica della meccanica quantistica, come si sa, coinvolge spazi a infinite dimensioni e i numeri complessi? Forse è proprio su questo versante che si può individuare più efficacemente una demarcazione significativa. Il modo più conveniente, seguendo la trattazione delineata in [1] al quale rimandiamo per un quadro più completo ed esauritivo, ci sembra quello di inquadrare la meccanica quantistica alla luce di quella classica.

Formulazione Lagrangiana-Hamiltoniana della meccanica

La dinamica classica è retta dall'equazione di Newton $F = ma$ per un punto materiale di massa m . Per semplicità ci limitiamo per ora a considerare una sola dimensione spaziale x . Per una forza conservativa F possiamo introdurre l'energia potenziale

$$U(x) - U(x_0) = - \int_{x_0}^x F(u) \, du.$$

Per il teorema fondamentale del calcolo integrale risulta $F(x) = -\frac{dU}{dx}$. La funzione Lagrangiana del sistema è introdotta considerando appunto l'energia potenziale e quella cinetica collegata alla velocità v della massa m . Precisamente

$$L(x, v) = \frac{1}{2}mv^2 - U(x).$$

L'equazione di Newton si può reinterpretare, piuttosto che mediante funzioni diverse, tramite la sola funzione Lagrangiana. Infatti, introducendo i simboli di derivate parziali per funzioni di più variabili, possiamo riscrivere

$$F(x) = -\frac{dU}{dx} = \frac{\partial L}{\partial x}.$$

D'altra parte, derivando rispetto a v

$$\frac{\partial L}{\partial v} = mv \Rightarrow m \frac{dv}{dt} = \frac{d}{dt} \left(\frac{\partial L}{\partial v} \right).$$

Siamo pronti per riformulare l'equazione di Newton in termini lagrangiani

$$F = ma \Leftrightarrow F = m \frac{dv}{dt} \Leftrightarrow \frac{\partial L}{\partial x} - \frac{d}{dt} \left(\frac{\partial L}{\partial v} \right) = 0.$$

Pertanto, l'equazione della dinamica è equivalente a un'equazione che coinvolge la sola funzione Lagrangiana, detta equazione di Eulero-Lagrange. Lo sforzo reinterpretativo è giustificato dalle notevoli proprietà matematiche connesse con le strutture collegate alla lagrangiana. Intanto, un fatto di fondamentale importanza, che l'equazione di Eulero-Lagrange corrisponde alle configurazioni stazionarie *dell'azione*

$$\mathcal{A} := \int_{t_0}^t L(x, v) dt.$$

Sovente, questa proprietà è condensata nel cosiddetto *Principio di minima azione* per cui l'evoluzione dinamica di un sistema è orientata dalla necessità di rendere minima (o massima) una azione (si veda [2] per un'introduzione a questi temi), cosa che suggerisce la possibilità di studiare la dinamica nel cosiddetto *spazio delle fasi* di coordinate (x, v) , studiando problemi di ottimizzazione (come il principio di Maupertis-Fermat nell'ottica). La corrispondente formulazione *duale* è dovuta ad Hamilton. L'hamiltoniana del sistema rappresenta l'energia totale data dall'energia cinetica più quella potenziale: $H = \frac{1}{2}mv^2 + U(x)$ (si tratta della cosiddetta trasformata di Fenchel-Legendre della lagrangiana). Se si introducono le cosiddette *coordinate generalizzate* tramite posizione e quantità di moto, ovvero $q = x$ e $p = mv$, le equazioni di Hamilton

$$\frac{dq}{dt} = \frac{\partial H}{\partial p}; \quad \frac{dp}{dt} = -\frac{\partial H}{\partial q} \quad (1)$$

caratterizzano completamente il sistema dinamico e la sua evoluzione nel tempo.

La generalità di questo approccio ne costituisce anche la forza. Gli enti rilevanti per la descrizione fisica diventano funzioni delle coordinate generalizzate. Diremo che *gli osservabili* della meccanica classica sono proprio le funzioni $f(q, p)$ sullo spazio delle fasi delle coordinate generalizzate. L'algebra delle funzioni fornisce automaticamente una struttura molto versatile e potente per l'analisi della dinamica. Una peculiarità decisiva è l'individuazione d un'algebra di Lie per gli osservabili. Si tratta di una relazione binaria, denotata usualmente con il simbolo $\{f, g\}$, che soddisfa le proprietà

1. (Linearità) $\{f, g + \alpha h\} = \{f, g\} + \alpha\{f, h\}$;
2. (Antisimmetria) $\{f, g\} = -\{g, f\}$;
3. (Identità di Jacobi) $\{f, \{g, h\}\} + \{g, \{h, f\}\} + \{h, \{f, g\}\} = 0$;
4. (Prodotto) $\{f, gh\} = g\{f, h\} + h\{f, g\}$.

L'ultima proprietà dovrebbe risultare familiare ricordando la formula di *derivazione del prodotto*. In effetti, nella meccanica classica una struttura di Algebra di Lie si realizza tramite le cosiddette *parentesi di Poisson*:

$$\{f, g\} := \frac{\partial f}{\partial p} \frac{\partial g}{\partial q} - \frac{\partial f}{\partial q} \frac{\partial g}{\partial p}. \quad (2)$$

Stati e osservabili

Anche la meccanica classica deve tener conto dell'incertezza e della probabilità. Una presenza sistematica di questi componenti si ha nella meccanica statistica. Il fatto è che non sempre il risultato di una misura sperimentale è individuato con esattezza. Se $f(q, p)$ è un osservabile relativo ad un certo esperimento, allora se una misurazione delle variabili (q_0, p_0) produce un risultato identificabile con certezza avremo un risultato $f(q_0, p_0)$ e diremo che lo stato è puro. Ma questa non è certamente la norma. Se indichiamo con w uno stato, ovvero l'insieme di tutte le condizioni che caratterizzano un certo set sperimentale, e con f gli osservabili relativi a quello stato, ovvero i risultati sperimentali determinabili nello stato w , più che valori univocamente determinati ci aspettiamo una certa distribuzione di probabilità. In altre parole, ad ogni osservabile f dello stato w è associata una $w_f \in \mathcal{P}(\mathbb{R})$ che descrive la distribuzione di probabilità relativa alle misure sperimentali. Gli stati allora possono essere *misti*, costituiti da una qualunque combinazione convessa di stati, vale a dire del tipo

$$w = \alpha w_1 + (1 - \alpha) w_2 \in \mathcal{P}(\mathbb{R})$$

per $w_1, w_2 \in \mathcal{P}(\mathbb{R})$ e $0 \leq \alpha \leq 1$. Diremo allora che uno stato è puro se non si può decomporre tramite combinazioni convesse, ovvero se $w = w_1 = w_2$. Se il risultato di una misura sperimentale è una probabilità, allora il suo valore rappresentativo è costituito dal *valor medio* (baricentro della misura di probabilità) definito da

$$(f|w) := \int_{\mathbb{R}} u dw_f(u). \quad (3)$$

Per gli stati puri risulta $w_f = \delta_{u_0}$ con δ la misura concentrata *delta di Dirac*; per ogni insieme misurabile $A \subset \mathbb{R}$ la delta di Dirac è definita da

$$\delta_u(A) = \begin{cases} 1 & u \in A \\ 0 & u \notin A \end{cases}.$$

Il baricentro della delta di Dirac è esattamente il centro della misura e $(f|w) = u_0$, per cui l'osservabile produce con certezza il valore $f(p_0, q_0) = u_0$.

La teoria richiede che gli stati soddisfino le basilari proprietà:

1. (Valor medio della costante) $(C|w) = C$;
2. (Linearità) $(f + \alpha g|w) = (f|w) + \alpha(g|h)$;
3. (Positività) $(f^2|w) \geq 0$.

Queste proprietà mettono in moto la *teoria della misura e dell'integrazione* (Teoremi di rappresentazione di Riesz) e la forma più generale per il valor medio di un osservabile può essere rappresentato direttamente sullo spazio delle fasi $\mathcal{M}$ tramite una misura $\mu_w \in \mathcal{P}(\mathcal{M})$ associata allo stato w tramite la seguente

$$(f|w) = \int_{\mathcal{M}} f(q, p) d\mu_w(q, p).$$

Gli stati puri sono esattamente quelli per cui la misura di rappresentazione è una delta di Dirac, ovvero $\mu_w = \delta_{(q_0, p_0)}$.

Una formulazione matematica della meccanica quantistica

Qualunque teoria fisica deve trattare di stati e osservabili. Che cosa allora distingue la meccanica classica da quella quantistica? Se indichiamo con lettere minuscole $a, b, c, \dots$ gli osservabili, il fatto di poter considerare strutture algebriche articolate richiede in realtà la calcolabilità di funzioni $f(a, b)$, per esempio per poter definire somme e prodotti di osservabili. L'osservazione cruciale è che la calcolabilità di una funzione di due variabili $f(a, b)$ richiede che gli osservabili a, b siano simultaneamente misurabili. Mentre nella meccanica classica questa questione non pone restrizione alcuna, nei fenomeni quantistici emergono osservabili intrinsecamente *non simultaneamente misurabili*, come le coordinate generalizzate posizione e momento, in accordo con le relazioni di indeterminazione di Heisenberg $\Delta x \Delta p \geq \frac{\hbar}{4\pi}$. La necessità di tener conto di questa simultaneità conduce direttamente al nocciolo del problema: chi sono allora gli osservabili della meccanica quantistica? Occorre pensare a qualcosa di nuovo giacché, qualunque cosa siano, gli osservabili non possono essere funzioni sullo spazio delle fasi. Inoltre, ammessa una nozione adeguata di osservabili, diciamo $A, B, C, \dots$ dovremo essere in grado di costruire un'algebra adeguata che al pari di quella classica consenta uno studio appropriato compatibile con le relazioni di indeterminazione. Una maniera naturale per passare dalla formulazione classica a

quella quantistica consiste nel passare dalle funzioni ai *funzionali*. Molto genericamente, un funzionale è una funzione $\mathcal{F} : X \rightarrow Y$ dove X, Y sono degli spazi vettoriali piuttosto generali (Spazi di Banach, di Hilbert, ecc.) anche di *dimensione infinita*. Esempi tipici provengono proprio dalla meccanica classica e dal *Calcolo delle Variazioni* (si veda anche [2]). Il principio di Fermat, per esempio, conduce a chiedersi quale sia il cammino più corto che connette due punti assegnati, producendo il cosiddetto *problema delle geodetiche*. Rispondere a questa domanda significa trovare configurazioni di minimo di un funzionale, quello che ad ogni curva dello spazio associa la sua lunghezza. In tal caso lo spazio X è costituito dallo spazio delle curve (di dimensione infinita) mentre $Y = \mathbb{R}$. Il funzionale in oggetto è in qualche modo una funzione su uno spazio di funzioni, una funzione di funzioni insomma, da cui il nome *funzionale*. Talvolta, specialmente nel caso lineare, i funzionali sono detti anche *operatori*. Lo studio sistematico di questi oggetti costituisce una branca molto vasta della matematica: l'Analisi Funzionale.

Matrici quantistiche

Per farsi un'idea di come l'Analisi Funzionale costituisca la struttura naturale per la meccanica quantistica, possiamo partire dal considerare la struttura più familiare delle matrici A, B, C quadrate di ordine $n \times n$, che rappresentano degli operatori lineari $v \in \mathbb{R}^n \mapsto Av \in \mathbb{R}^n$. Naturalmente, se le nostre matrici rappresentano gli osservabili della teoria, il problema è definire in maniera sensata delle funzioni $f(A)$ sulle matrici. A tal fine, si può procedere considerando le *matrici autoaggiunte* e il cosiddetto *Teorema Spettrale*. Per capire di cosa si tratta, ricordiamo dapprima la nozione di prodotto scalare in $\mathbb{R}^n$. In coordinate cartesiane: se $u = (u_1, u_2, \dots, u_n)$ e $v = (v_1, v_2, \dots, v_n)$ sono due vettori di $\mathbb{R}^n$, allora il prodotto scalare è definito da

$$\langle u, v \rangle = \sum_{i=1}^n u_i v_i.$$

Ricordiamo inoltre che una matrice quadrata A di ordine n è legata alla sua trasposta A^t (che si ottiene scambiando tra loro righe e colonne) dalla fondamentale relazione

$$\forall \mathbf{u}, \mathbf{v} \in \mathbb{R}^n, \langle A\mathbf{u}, \mathbf{v} \rangle = \langle \mathbf{u}, A^t \mathbf{v} \rangle. \quad (4)$$

Infatti, se $A = (a_{ij})$, per definizione di prodotto *righe per colonne* abbiamo che

$$\begin{aligned} \langle A\mathbf{u}, \mathbf{v} \rangle &= \sum_{j=1}^n \left(\sum_{i=1}^n a_{ji} u_i \right) v_j = \sum_{j=1}^n \left(\sum_{i=1}^n a_{ji} u_i v_j \right) = \sum_{i=1}^n \left(\sum_{j=1}^n a_{ji} u_i v_j \right) = \\ &= \sum_{i=1}^n \left(\sum_{j=1}^n a_{ji} v_j \right) u_i = \sum_{i=1}^n (A^t \mathbf{v})_i u_i = \langle \mathbf{u}, A^t \mathbf{v} \rangle. \end{aligned}$$

La matrice aggiunta A^* di A è definita proprio tramite la matrice trasposta, ovvero $A^* := A^t$. Diremo che una matrice è autoaggiunta se essa coincide con la propria aggiunta, vale a dire se $A = A^*$. Il menzionato Teorema Spettrale garantisce che ogni matrice autoaggiunta possiede n autovalori $\lambda_i \in \mathbb{R}$ e n autovettori $v_i \in \mathbb{R}^n$ che soddisfano la relazione

$$Av_i = \lambda_i v_i, \quad i = 1, \dots, n. \quad (5)$$

L'insieme degli autovalori costituisce lo *spettro* della matrice A . Inoltre, gli autovettori v_i corrispondenti ad autovalori distinti sono ortogonali e costituiscono quindi una base dello spazio. Questi risultati forniscono un modo del tutto naturale per definire delle funzioni sulle matrici, giacché le matrici autoaggiunte restano univocamente determinate dal loro spettro. Se f è una funzione reale, basterà definire $f(A)$ come la matrice autoaggiunta avente per spettro le immagini $f(\lambda_i)$. Il ricorso alle matrici autoaggiunte richiede però qualche sforzo per determinare la struttura algebrica idonea. Intanto valutiamo per esercizio che

$$(\alpha A)^* = \alpha A^*; \quad (AB)^* = B^* A^*.$$

Infatti, utilizzando la relazione tra matrici e prodotto scalare abbiamo

$$\langle \alpha A \mathbf{u}, \mathbf{v} \rangle = \alpha \langle \mathbf{u}, A^* \mathbf{v} \rangle = \langle \mathbf{u}, \alpha A^* \mathbf{v} \rangle \Rightarrow (\alpha A)^* = \alpha A^*.$$

Lasciamo al lettore il compito di verificare che inoltre $(A + B)^* = A^* + B^*$.

$$\langle AB \mathbf{u}, \mathbf{v} \rangle = \langle B \mathbf{u}, A^* \mathbf{v} \rangle = \langle \mathbf{u}, B^* A^* \mathbf{v} \rangle \Rightarrow (AB)^* = B^* A^*.$$

L'ultima relazione trovata ci dice che il prodotto righe per colonne tra le matrici non è la nozione giusta di prodotto, giacché il prodotto AB di due matrici autoaggiunte non è autoaggiunto. Si definisce allora il prodotto simmetrico

$$A \circ B := \frac{AB + BA}{2}.$$

Si verifica facilmente che il prodotto simmetrico è autoaggiunto:

$$(A \circ B)^* = \frac{1}{2} ((AB)^* + (BA)^*) = \frac{1}{2} (B^* A^* + A^* B^*) = \frac{1}{2} (BA + AB) = A \circ B.$$

Per le matrici, un operatore fondamentale è la *traccia*, denotata con $Tr(A)$, definita dalla somma degli autovalori:

$$Tr(A) := \sum_{i=1}^n a_{ii} = \sum_{i=1}^n \langle A e_i, e_i \rangle = \sum_{i=1}^n \lambda_i \quad (6)$$

dove $e_1 = (1, 0, \dots, 0)$, $e_2 = (0, 1, \dots, 0)$ etc rappresentano i vettori della base canonica di $\mathbb{R}^n$. Si può verificare (il lettore può svolgerlo quale utile esercizio) che la traccia non dipende dalla scelta della base e inoltre che $Tr(AB) = Tr(BA)$ quali che siano le matrici A, B .

I teoremi di rappresentazione per stati e osservabili nel caso delle matrici suggeriscono di assegnare una matrice autoaggiunta M tale che $\langle Mu, u \rangle \geq 0$ (definita positiva) con $Tr(M) = 1$, detta matrice densità. Il valor medio assume allora la seguente forma

$$(A|w) := Tr(MA) = \sum_{i=1}^n \mu_i \langle A\varphi_i, \varphi_i \rangle,$$

dove μ_i sono gli autovalori di M e φ_i i corrispondenti autovettori. Per ottenere la seconda uguaglianza si può assumere che la base di autovettori sia quella canonica dello spazio (altrimenti basta operare un cambio di base) per cui

$$Tr(MA) = \sum_{i=1}^n \langle MA\varphi_i, \varphi_i \rangle = \sum_{i=1}^n \langle A\varphi_i, M\varphi_i \rangle = \sum_{i=1}^n \mu_i \langle A\varphi_i, \varphi_i \rangle. \quad (7)$$

Pertanto, analogamente al caso classico, si può interpretare lo stato misto come la sovrapposizione di stati puri P_{φ_i} ciascuno con probabilità μ_i , dove lo stato puro è determinato dalla proiezione (prodotto scalare) lungo l'autovettore. Precisamente, se ψ è un autovettore, lo stato puro corrispondente è definito dall'operatore $P_{\psi}(\eta) = \langle \psi, \eta \rangle \eta$. Per uno stato puro il valor medio di un'osservabile assume la forma

$$(A|w) = \langle A\psi, \psi \rangle. \quad (8)$$

Una interpretazione analoga può essere data per gli autovalori e gli autovettori dell'osservabile A (si veda [1, cap. 8]): gli autovalori λ_i rappresentano i possibili risultati delle misurazioni sull'osservabile A nello stato w con probabilità $\langle Mv_i, v_i \rangle$, dove v_i sono i corrispondenti autovettori. Se il sistema si trova in uno stato puro corrispondente a un autovettore v_i di A , allora la misura λ_i è ottenuta con certezza.

Per chiudere il cerchio è desiderabile costruire un'algebra di Lie analoga a quanto accade per la meccanica classica. Il candidato naturale è il cosiddetto *commutatore*

$$[A, B] = AB - BA.$$

Andiamo a valutare la sussistenza delle quattro proprietà di Lie:

1. Linearità: $[A, B + \alpha C] = A(B + \alpha C) - (B + \alpha C)A = AB + \alpha AC - BA - \alpha CA = [A, B] + \alpha[A, C]$;
2. Antisimmetria: $[A, B] = AB - BA = -(BA - AB) = -[B, A]$;
3. Prodotto: $[A, B \circ C] = [A, \frac{BC+CB}{2}] = \frac{1}{2}(ABC + ACB - BCA - CBA) = \frac{1}{2}B(AC - CA) - \frac{1}{2}BAC + \frac{1}{2}(AC - CA)B + \frac{1}{2}CAB + \frac{1}{2}ABC - \frac{1}{2}CBA = B \circ [A, C] + \frac{1}{2}C(AB - BA) + \frac{1}{2}(AB - BA)C = B \circ [A, C] + C \circ [A, B]$;
4. Jacobi: $[A, [B, C]] + [B, [C, A]] + [C, [A, B]] = [A, BC - CB] + [B, (CA - AC)] + [C, (AB - BA)] = ABC - ACB - BCA + CBA + BCA - BAC - CAB + ACB + CAB - CBA - ABC + BAC = 0$;

Sembra funzionare tutto a dovere ma tutti sanno che nella meccanica quantistica sono coinvolti i numeri complessi. Come mai? Il commutatore ha tutte le proprietà desiderate ma abbiamo trascurato di controllare se esso produce un operatore autoaggiunto. In effetti

$$[A, B]^* = (AB - BA)^* = B^*A^* - A^*B^* = BA - AB = -[A, B].$$

Un segno meno guastafeste ci separa da un'algebra efficiente per la meccanica quantistica. Ma la matematica ci soccorre facilmente a patto di passare dal campo reale $\mathbb{R}$ al campo complesso $\mathbb{C}$. Pochi accorgimenti permettono di dipanare la situazione. Ricordando l'operazione di *coniugazione* $z = \alpha + i\beta \mapsto \bar{z} = \alpha - i\beta$, il prodotto scalare tra vettori nel campo complesso si definisce tramite la seguente formula:

$$\langle x, y \rangle := \sum_{i=1}^n x_i \bar{y}_i,$$

da cui segue che $\langle x, x \rangle = |x|^2$ è un numero reale positivo. Per la nozione di operatori aggiunti, si definisce invece $A^* := \bar{A}^t$ da cui segue che $(\alpha A)^* = \bar{\alpha} A^*$. Il Teorema Spettrale continua a valere e lo spettro resta composto da autovalori tutti reali, salvando l'esigenza di misurazioni sperimentali costituite da numeri reali. Per ottenere l'algebra di Lie basta porre $[A, B] = ki(AB - BA)$. La verifica che si tratti di un operatore autoaggiunto è quasi superflua:

$$[A, B]^* = -ki(AB - BA)^* = -ki(BA - AB) = ki(AB - BA) = [A, B].$$

La costante k va determinata in modo che la teoria si adatti ai dati sperimentali. Essi mostrano che $k = \frac{1}{\hbar}$, dove $\hbar$ è la costante di Planck (ridotta). Possiamo ora definire le parentesi di Poisson quantistiche

$$\{A, B\}_h := \frac{i}{\hbar} (AB - BA).$$

Pertanto, la necessità di lavorare in campo complesso serve a garantire l'applicazione in tutte le circostanze della teoria degli operatori autoaggiunti. Per quanto bizzarra possa risultare la meccanica quantistica, i risultati delle misurazioni devono comunque essere numeri reali. Verifichiamo, come già accennato, che gli autovalori di un operatore autoaggiunto sono reali:

$$\lambda = \lambda \langle v, v \rangle = \langle Av, v \rangle = \langle v, Av \rangle = \bar{\lambda} \langle v, v \rangle = \bar{\lambda} \Leftrightarrow \lambda \in \mathbb{R} \quad (9)$$

essendo l'autovettore v relativo all'autovalore λ di modulo unitario. Come osservato, il prodotto di autoaggiunti non è necessariamente autoaggiunto. Tuttavia, l'operatore traccia restituisce naturalmente un numero reale. Questa circostanza ci assicura che il valor medio di un operatore autoaggiunto è ancora un numero reale. A tal fine basta richiamare le proprietà di base del coniugio complesso:

$$\overline{\langle x, y \rangle} = \overline{\sum_{i=1}^n x_i \bar{y}_i} = \sum_{i=1}^n \bar{x}_i y_i = \sum_{i=1}^n \bar{x}_i y_i = \langle y, x \rangle.$$

Pertanto, per una matrice autoaggiunta, quale che sia il vettore u abbiamo

$$\overline{\langle Au, u \rangle} = \langle u, Au \rangle = \langle Au, u \rangle \Leftrightarrow \langle Au, u \rangle \in \mathbb{R}.$$

Anche in campo complesso ha senso considerare matrici definite positive e la relativa nozione di matrice densità. Ricordando l'espressione del valor medio

$$(A|w) = \text{Tr}(MA) = \sum_{i=1}^n \mu_i \langle A\varphi_i, \varphi_i \rangle \in \mathbb{R}. \quad (10)$$

Visto che siamo in tema, verifichiamo le altre proprietà del valor medio per osservabili quantistici. Per la proprietà di positività basta osservare che gli autovalori $\mu_i = \langle M\varphi_i, \varphi_i \rangle \geq 0$. Allora il valor medio è un funzionale positivo:

$$(A^2|w) = \text{Tr}(A^2M) = \sum_{i=1}^n \mu_i \langle A^2\varphi_i, \varphi_i \rangle = \sum_{i=1}^n \mu_i \langle A\varphi_i, A\varphi_i \rangle = \sum_{i=1}^n \mu_i |A\varphi_i|^2 \geq 0.$$

La linearità è una diretta conseguenza della linearità del prodotto scalare

$$\begin{aligned} (A + \alpha B|w) &= \text{Tr}(M(A + \alpha B)) = \sum_{i=1}^n \mu_i \langle (A + \alpha B)\varphi_i, \varphi_i \rangle = \\ &= \sum_{i=1}^n \mu_i \langle A\varphi_i, \varphi_i \rangle + \alpha \sum_{i=1}^n \mu_i \langle B\varphi_i, \varphi_i \rangle = (A|w) + \alpha(B|w). \end{aligned}$$

Infine, per quanto riguarda il valor medio delle costanti, data la linearità, è sufficiente controllare il comportamento del valor medio rispetto alla matrice identità I :

$$(I|w) = \text{Tr}(MI) = \text{Tr}(M) = 1.$$

Il Teorema spettrale

Per garantire un minimo di completezza, diamo qualche dettaglio su questo fondamentale risultato. Abbiamo già appurato che gli autovalori associati ad una matrice (operatore) autoaggiunta sono tutti reali. Ma chi ci garantisce che esistano sempre? Dire che v è un autovettore significa cercare una soluzione, diversa dal vettore nullo, dell'equazione $Av = \lambda v$. Dividendo per il modulo $|v|$ potremo considerare vettori unitari. Equivalentemente, detta I la matrice identità, possiamo considerare l'equazione

$$(A - \lambda I)v = 0$$

che corrisponde a un sistema lineare di n equazioni. L'algebra lineare permette di studiare agevolmente l'esistenza di soluzioni. Se $\det(A - \lambda I) \neq 0$ allora il sistema ammette una ed una sola soluzione, che pertanto non può che essere il vettore nullo $v = 0$. Ma noi siamo interessati a soluzioni diverse da zero. Perveniamo così alla cosiddetta *equazione caratteristica*:

$$\det(A - \lambda I) = 0. \quad (11)$$

Dunque, λ è un autovalore di A se e soltanto se λ è una radice di (11). Inoltre, la risoluzione del sistema lineare fornisce il corrispondente autovettore. Ora, la (11) è un'equazione, nel campo complesso, di grado n nella variabile λ . Il teorema fondamentale dell'algebra (si veda anche [3] assicura l'esistenza di esattamente n radici complesse (contate con la rispettiva molteplicità). Resta da verificare il fatto che gli autovettori formano una base. Verifichiamo che ad autovalori distinti $\lambda_i \neq \lambda_j$ corrispondono autovettori v_i, v_j ortogonali tra loro:

$$\lambda_i \langle v_i, v_j \rangle = \langle Av_i, v_j \rangle = \langle v_i, Av_j \rangle = \lambda_j \langle v_i, v_j \rangle \Leftrightarrow (\lambda_i - \lambda_j) \langle v_i, v_j \rangle = 0 \Leftrightarrow \langle v_i, v_j \rangle = 0.$$

Indeterminazione di Heisenberg

Lo sforzo per sviluppare una matematica capace di descrivere i fenomeni quantistici non si deve naturalmente infrangere sulle relazioni di indeterminazioni di Heisenberg che hanno motivato la scelta degli operatori autoaggiunti come osservabili della meccanica quantistica. In effetti, le relazioni di indeterminazione trovano in questa formulazione uno sviluppo naturale. L'incertezza su un osservabile è statisticamente espressa dalla varianza. Fissato lo stato w di interesse, indicato con $A_m := (A|w)$ il valore medio, la varianza, che indichiamo col simbolo $\Delta_w A$, è legata al valor medio del quadrato della differenza tra l'osservabile e il suo valor medio. Precisamente la varianza è definita da

$$\Delta_w^2 A := ((A - A_m)^2 | w). \quad (12)$$

Consideriamo per semplicità uno stato puro (con qualche accorgimento in più si possono ugualmente trattare gli stati sovrapposti) di autovettore di modulo unitario ψ . Per un arbitrario $\alpha \in \mathbb{R}$, partiamo dalla positività del prodotto scalare

$$\langle (A + i\alpha B)\psi, (A + i\alpha B)\psi \rangle \geq 0. \quad (13)$$

Sviluppando il prodotto scalare, ricordando che A, B sono autoaggiunti, si ottiene

$$\begin{aligned} \langle A\psi, A\psi \rangle + \langle A\psi, i\alpha B\psi \rangle + \langle i\alpha B\psi, A\psi \rangle + \langle i\alpha B\psi, i\alpha B\psi \rangle = \\ \langle A^2\psi, \psi \rangle - i\alpha \langle BA\psi, \psi \rangle + i\alpha \langle AB\psi, \psi \rangle + \alpha^2 \langle B^2\psi, \psi \rangle = \\ \langle A^2\psi, \psi \rangle + i\alpha \langle (AB - BA)\psi, \psi \rangle + \alpha^2 \langle B^2\psi, \psi \rangle. \end{aligned}$$

Trattandosi di uno stato puro, ricordando la (8), dalla (13) si ottiene

$$(A^2|w) + \alpha^2(B^2|w) + \alpha\hbar(\{A, B\}_h|w) \geq 0.$$

La disuguaglianza appena ottenuta è valida quale che sia $\alpha \in \mathbb{R}$. Il valor medio è un numero reale e in particolare $(B^2|w) \geq 0$. Il trinomio di secondo grado nella variabile α deve pertanto avere un discriminante negativo:

$$\hbar^2 (\{A, B\}_h|w)^2 - 4(A^2|w)(B^2|w) \leq 0 \Leftrightarrow (A^2|w)(B^2|w) \geq \frac{\hbar^2}{4} (\{A, B\}_h|w)^2. \quad (14)$$

Per ottenere la varianza occorre sostituire nella formula precedente $A \mapsto A - A_M$, $B \mapsto B - B_M$, dove $A_M = (A|w)$, $B_M = (B|w)$. Valutiamo preliminarmente che le parentesi di Poisson restano invariate

$$\begin{aligned}\{A - A_M, B - B_M\}_h &= \frac{i}{\hbar} ((A - A_M)(B - B_M) - (B - B_M)(A - A_M)) = \\ &= \frac{i}{\hbar} (AB - B_M A - A_M B + A_M B_M - BA + A_M B + B_M A - A_M B_M) = \\ &= \frac{i}{\hbar} (AB - BA) = \{A, B\}_h.\end{aligned}$$

Pertanto, dalla (14) otteniamo la relazione di indeterminazione di Heisenberg

$$\Delta_w A \Delta_w B \geq \frac{\hbar}{2} |(\{A, B\}_h|w)|. \quad (15)$$

In generale, la varianza di due osservabili non è simultaneamente nulla e i due osservabili non sono pertanto simultaneamente misurabili. Questo fenomeno, come si vede dalla (15), è strettamente collegato alla commutatività degli operatori. Il fatto che in meccanica quantistica le cose si complichino alquanto discende proprio dal fatto di lavorare su un'algebra non commutativa. Questa circostanza costituisce una netta differenza con la meccanica classica. Mentre in quest'ultima possiamo immaginare di raffinare sempre più la precisione dello stato w in modo da rendere le incertezze sugli osservabili f, g , almeno in linea di principio, arbitrariamente piccole, in meccanica quantistica, se gli osservabili A, B non commutano allora non c'è niente da fare: la (15) pone una limitazione ineludibile sulle incertezze.

Posizione e quantità di moto

Quanto detto fino ad ora può risultare piuttosto astratto, in fin dei conti è la rispondenza ai dati sperimentali a costituire il principio decisivo per qualunque teoria. Naturalmente, quale che sia la formulazione per la meccanica quantistica, è desiderabile che essa risulti *compatibile* con quella classica. Vorremmo cioè una corrispondenza tra un osservabile classico f e un corrispettivo osservabile quantistico A_f , in modo per esempio che le predizioni della meccanica quantistica su una situazione macroscopica sia confrontabile con la descrizione classica. La corrispondenza $f \mapsto A_f$ non sarà biunivoca, giacché vogliamo contemplare situazioni quantistiche che non abbiano un corrispettivo classico. Ma certamente la corrispondenza $f \mapsto A_f$ renderà la teoria confrontabile con i dati sperimentali. La meccanica quantistica da questo punto di vista è ineguagliabile fornendo una corrispondenza con le misurazioni sperimentali di altissima precisione. Inoltre, la corrispondenza $f \mapsto A_f$ consente di riguardare la meccanica classica come limite per $\hbar \rightarrow 0$ di quella quantistica (si veda [1, cap. 14]).

Per farci un'idea di come questa corrispondenza possa avvenire consideriamo le coordinate generalizzate (q, p) , posizione e quantità di moto. Vogliamo trovare i corrispondenti osservabili quantistici $q \mapsto Q, p \mapsto P$. Per identificare questi corrispondenti osservabili quantistici, che chiameremo ancora posizione e quantità di

moto, calcoliamo le parentesi di Poisson classiche. Nelle tre coordinate spaziali, per la posizione dobbiamo calcolare $\{q_i, q_j\}$ per gli osservabili $f(q, p) = q_i$, $g(q, p) = q_j$, con $i, j = 1, 2, 3$. Calcolando le derivate si ottiene

$$\{q_i, q_j\} = \frac{\partial f}{\partial p} \frac{\partial g}{\partial q} - \frac{\partial f}{\partial q} \frac{\partial g}{\partial p} = 0.$$

Similmente, valutiamo

$$\{p_i, p_j\} = 0; \quad \{p_i, q_j\} = \delta_{i,j}$$

dove la delta di Kronecker produce $\delta_{i,j} = 1$ per indici uguali $i = j$. Mentre $\delta_{i,j} = 0$ per $i \neq j$. L'idea è quella di determinare degli osservabili quantistici che producano lo stesso comportamento per le parentesi di Poisson, ovvero che soddisfano le cosiddette *relazioni di quantizzazione-commutazione di Heisenberg*

$$\{Q_i, Q_j\}_h = 0; \quad \{P_i, P_j\}_h = 0; \quad \{P_i, Q_j\}_h = \delta_{ij}. \quad (16)$$

Per i corrispondenti quantistici di posizione e quantità di moto, la relazione di indeterminazione di Heisenberg, per lo stato relativo al moto di una singola particella, assume l'usuale forma

$$\Delta_w P \Delta_w Q \geq \frac{\hbar}{2} |(\{P, Q\}_h|w)| = \frac{\hbar}{2} |(I|w)| = \frac{\hbar}{2}.$$

Le dimensioni contano

Il modello che abbiamo sviluppato utilizzando le matrici si rivela troppo riduttivo per gli scopi minimi richiesti dalla fisica. Sia per motivi applicativi che teorici. Dal punto di vista degli esperimenti, l'algebra delle matrici ha il difetto che ogni osservabile può essere collegato con al più n possibili risultati reali corrispondenti agli autovalori. Cosa che nel misurare una grandezza fisica sarebbe una restrizione troppo severa. Nella maggior parte dei casi vorremmo addirittura un intervallo continuo di valori possibili per una certa misurazione. Dal punto di vista teorico, per esempio, chi garantisce che esistano gli osservabili quantistici che soddisfano le relazioni di Heisenberg (16)? In effetti, non è difficile mostrare che in dimensione finita tali osservabili non si possono trovare. In effetti, ci aspettiamo che posizione e quantità di moto debbano corrispondere ad operatori non limitati, cosa che le matrici non possono produrre. Come accennato in precedenza, gli strumenti dell'analisi funzionale permettono di trattare il caso di dimensioni infinite con molta efficacia. Gli spazi dotati di prodotto scalare di dimensione infinita conducono a considerare la nozione di spazio di Hilbert. Lo spettro degli operatori negli spazi di Hilbert è ancora composto da numeri reali e può essere costituito da infiniti elementi o anche da intervalli. Il Teorema di von Neumann-Stone assicura inoltre l'esistenza degli osservabili quantistici (operatori autoaggiunti in uno spazio di Hilbert) che soddisfano relazioni come quelle di Heisenberg (16).

Anche l'operatore traccia si definisce in dimensione infinita come una serie

$$Tr(A) := \sum_{i=1}^{+\infty} \langle Av_i, v_i \rangle = \sum_{i=1}^{+\infty} \lambda_i$$

dando luogo a possibili valori non limitati. Resta valido anche in dimensione infinita il fatto che la traccia non dipende dalla scelta della base e che, sebbene due operatori possano anche non commutare, $Tr(AB) = Tr(BA)$.

I teoremi di rappresentazione dell'analisi funzionale consentono di formulare questi concetti tramite integrazione in opportuni spazi di funzioni. Si tratta degli spazi di Lebesgue $L^2(\mathbb{R}^3)$ costituiti dalle funzioni su $\mathbb{R}^3$, a valori nel campo complesso, per cui il quadrato del modulo ammette integrale di Lebesgue finito. Il prodotto scalare in questo spazio è definito da

$$\langle f, g \rangle := \int_{\mathbb{R}^3} f(x) \overline{g(x)} dx.$$

Questa veste più *esplicita* consente di dare una apparenza più *concreta* alla formulazione quantistica. Per esempio, gli operatori quantistici per posizione e quantità di moto possono essere definiti da

$$\forall \varphi \in L^2(\mathbb{R}^3) : Q_j \varphi(x) = x_j \varphi(x); \quad P_j \varphi(x) = \frac{\hbar}{i} \frac{\partial \varphi}{\partial x_j}(x).$$

Il lettore attento noterà che questi osservabili non sono limitati e che la loro definizione richiede una restrizione sul dominio, chi assicura infatti che la derivata di una funzione in L^2 sia ancora in L^2 ? Queste considerazioni conducono allo studio di spazi di funzioni particolari come quelli di Sobolev, spazi a momento finito ecc.

Verifichiamo le relazioni di commutazione di Heisenberg. Non è restrittivo, per approssimazione delle funzioni integrabili con funzioni regolari, considerare una funzione test $\varphi \in \mathcal{C}_c^\infty(\mathbb{R}^3)$ infinite volte derivabile e nulla al di fuori di un insieme compatto (supporto compatto).

$$\{Q_i, Q_j\}_h \varphi = \frac{i}{\hbar} (Q_i(Q_j \varphi) - Q_j(Q_i \varphi)) = \frac{i}{\hbar} (x_i x_j \varphi(x) - x_j x_i \varphi(x)) = 0$$

visto la commutatività del prodotto tra numeri reali. Similmente per i momenti

$$\{P_i, P_j\}_h \varphi = \frac{i}{\hbar} (P_i(P_j \varphi) - P_j(P_i \varphi)) = \left(\frac{\partial^2 \varphi}{\partial x_i \partial x_j}(x) - \frac{\partial^2 \varphi}{\partial x_j \partial x_i}(x) \right) = 0.$$

Questa volta i due contributi si elidono a vicenda perché, per funzioni regolari, l'ordine di derivazione non conta (Teorema di Schwarz). Infine

$$\begin{aligned} \{P_i, Q_j\}_h \varphi &= \frac{i}{\hbar} (P_i(Q_j \varphi) - Q_j(P_i \varphi)) = \left(\frac{\partial}{\partial x_i} (x_j \varphi(x)) - x_j \frac{\partial \varphi}{\partial x_i}(x) \right) = \\ &= x_j \frac{\partial \varphi}{\partial x_i}(x) - x_j \frac{\partial \varphi}{\partial x_i}(x) = 0. \end{aligned}$$

L'equazione di Schrodinger

Quanto espresso sino ad ora è in un certo senso una forma di cinematica. Per la dinamica occorre prescrivere l'evoluzione nel tempo. Limitandoci ad una particella in

una sola dimensione spaziale per semplicità, il suo stato è individuato dagli autovettori, individuati da una funzione $f \in L^2(\mathbb{R})$. Se, conformemente al dualismo onda-particella, interpretiamo l'evoluzione in termini di un'onda progressiva, allora gli stati puri sono individuati da una *funzione d'onda*

$$\psi(x, t) = f(kx - wt)$$

dove k è il numero d'onda e w la frequenza angolare. Nel caso di un'onda armonica semplice (non è restrittivo sviluppando in serie di Fourier), utilizzando le formule di Eulero si può considerare

$$\psi(x, t) = e^{i(kx - wt)}.$$

Derivando si ottiene

$$\frac{\partial \psi}{\partial t}(x, t) = -wi\psi; \quad \frac{\partial \psi}{\partial x}(x, t) = ki\psi.$$

D'altra parte, ricordando che per le onde $w = 2\pi\nu$ e $k = \frac{2\pi}{\lambda}$, tenuto conto delle condizioni di quantizzazione di Planck $E = h\nu$ e $p = \frac{h}{\lambda}$, possiamo riscrivere le derivate della funzione d'onda come

$$i\hbar \frac{\partial \psi}{\partial t}(x, t) = E\psi; \quad -i\hbar \frac{\partial \psi}{\partial x}(x, t) = p\psi. \quad (17)$$

Le relazioni (17) suggeriscono di identificare gli operatori $i\hbar \frac{\partial}{\partial t}$ e $-i\hbar \frac{\partial}{\partial x}$ rispettivamente con gli operatori lineari energia E e quantità di moto p . D'altra parte, l'energia classica è fornita da

$$E = \frac{1}{2}mv^2 + U(x) = \frac{1}{2m}p^2 + U(x).$$

Dalle (17) otteniamo

$$i\hbar \frac{\partial \psi}{\partial t}(x, t) = E\psi = \frac{1}{2m}(i\hbar)^2 \left(\frac{\partial \psi}{\partial x}(x, t) \right)^2 + U(x)\psi.$$

Ponendo l'hamiltoniano $H := -\frac{\hbar^2}{2m} \left(\frac{\partial}{\partial x} \right)^2 + U$, si ottiene l'equazione di Schrodinger

$$i\hbar \frac{\partial \psi}{\partial t}(x, t) = H\psi.$$

Nell'ambito degli spazi di Hilbert occorrerà stabilire a seconda del contesto il significato di soluzione dell'equazione di Schrodinger: soluzioni deboli, distribuzioni, ecc. (si veda anche [4]).

Dalla dinamica classica a quella quantistica

Anche la dinamica quantistica discende naturalmente dal formalismo lagrangiano-hamiltoniano. Supponendo che l'hamiltoniana H del sistema soddisfi le condizioni

di regolarità per produrre esistenza e unicità delle soluzioni delle equazioni (1), resta definito il *flusso hamiltoniano*

$$G_t : \mathcal{M} \rightarrow \mathcal{M}, \quad G_t(q_0, p_0) = (q(q_0, p_0, t), p(q_0, p_0, t))$$

che ad ogni condizione iniziale (q_0, p_0) associa l'unica soluzione delle equazioni di Hamilton. Esistenza e unicità delle soluzioni si traducono nel fatto che il flusso hamiltoniano soddisfa le relazioni

$$G_t \circ G_s = G_{t+s}; \quad G_t^{-1} = G_{-t}. \quad (18)$$

La descrizione hamiltoniana della dinamica consiste nel seguire l'evoluzione di un osservabile $f(q, p)$ lungo il flusso hamiltoniano, vale a dire considerando un osservabile in funzione del tempo definito da

$$f_t(q, p) = f(G_t(q, p)) = f(q(q, p, t), p(q, p, t)). \quad (19)$$

L'equazione delle dinamica si ottiene considerando $\varphi(s) := f_t(G_s(q, p))$ e calcolando $\varphi'(0)$. Dalla derivazione della funzione composta abbiamo

$$\varphi'(s) = \langle \nabla f_t(G_s(q, p)), \frac{d}{ds} G_s(q, p) \rangle. \quad (20)$$

Ora, per definizione di flusso hamiltoniano abbiamo che $G_0(q, p) = (q, p)$ mentre

$$\frac{d}{ds} G_s(q, p) = \left(\frac{\partial H}{\partial p}(q, p), -\frac{\partial H}{\partial q}(q, p) \right).$$

Dalla (20), ricordando l'espressione delle parentesi di Poisson (2), si ottiene

$$\varphi'(0) = \frac{\partial f_t}{\partial q}(q, p) \frac{\partial H}{\partial p}(q, p) - \frac{\partial f_t}{\partial p}(q, p) \frac{\partial H}{\partial q}(q, p) = \{H, f_t\}. \quad (21)$$

D'altra parte, utilizzando le proprietà del flusso (18)

$$\varphi(s) = f_t(G_s(q, p)) = f(G_t(G_s(q, p))) = f(G_{t+s}(q, p)) = f(q(q, p, t+s), p(q, p, t+s)) \Rightarrow$$

$$\varphi'(0) = \frac{\partial f_t}{\partial t}(q, p).$$

La (21) produce allora l'equazione di evoluzione di Hamilton

$$\frac{\partial f_t}{\partial t} = \{H, f_t\}. \quad (22)$$

Il punto di vista hamiltoniano contempla pertanto uno stato fisso w in cui evolvono gli osservabili f_t . Con una metafora teatrale, potremmo dire che gli attori (osservabili) si muovono su un palcoscenico fisso. Ma, in modo forse meno intuitivo, si potrebbero anche tener fissi gli attori e far muovere il palcoscenico.

Questo secondo contesto corrisponde alla cosiddetta *formulazione di Liouville* della meccanica, che corrisponde esattamente all'equazione di Schrodinger nella quale a evolvere sono proprio gli stati. La formulazione classica di Hamilton corrisponde invece alla formulazione della dinamica quantistica di Heisenberg. Le due formulazioni sono equivalenti. Nel caso classico, consideriamo la misura di rappresentazione sullo spazio delle fasi che per semplicità consideriamo essere assolutamente continua con densità ρ . In altre parole, consideriamo un valor medio della forma

$$(f_t|w) = \int_{\mathcal{M}} f_t(u)\rho(u) du = \int_{\mathcal{M}} f(G_t(u))\rho(u) du.$$

Effettuando il cambio di variabili $v = G_t(u)$ otteniamo

$$(f_t|w) = \int_{\mathcal{M}} f(G_t(u))\rho(u) du = \int_{\mathcal{M}} f(v)\rho(G_t^{-1}(v))J(v) dv$$

dove $J(v)$ è lo Jacobiano del cambio di variabili (in questo caso specifico, in conseguenza del Teorema di Liouville $J(v) = 1$, si veda [1, cap. 3]). Considerando ora l'osservabile f fisso e la misura variabile di densità $\rho_t(v) = \rho(G_t^{-1}(v))J(v)$, abbiamo uno stato variabile w_t per cui le due descrizioni della dinamica coincidono

$$(f_t|w) = (f|w_t).$$

Ripercorrendo quanto presentato per la (22), si può verificare che l'equazione di evoluzione di Liouville assuma la forma

$$\frac{\partial \rho_t}{\partial t} = -\{H, \rho_t\}.$$

L'equazione di Heisenberg è la riformulazione quantistica della (22) in termini di operatori. Se M è l'operatore densità associato allo stato w , la descrizione di Heisenberg assume la forma

$$\frac{dM}{dt} = 0; \quad \frac{dA(t)}{dt} = \{H, A(t)\}_h. \quad (23)$$

L'equazione (23) di Heisenberg si può risolvere esplicitamente introducendo l'operatore esponenziale

$$e^A := \sum_{n=0}^{+\infty} \frac{A^n}{n!}$$

che formalmente corrisponde all'espansione in serie dell'usuale funzione esponenziale. Se $A = A(0)$ è il dato iniziale, la soluzione di (23) assume la forma

$$A(t) = e^{\frac{i}{\hbar}Ht} A e^{-\frac{i}{\hbar}Ht}. \quad (24)$$

In generale, gli esponenziali nella (24) non si semplificano giacché gli operatori possono non commutare tra loro. Se A e H commutano si ottengono invece soluzioni stazionarie. Effettuando la derivata della (24) si ottiene infatti

$$A'(t) = \frac{i}{\hbar} H e^{\frac{i}{\hbar}Ht} A e^{-\frac{i}{\hbar}Ht} - \frac{i}{\hbar} e^{\frac{i}{\hbar}Ht} A e^{-\frac{i}{\hbar}Ht} H = \frac{i}{\hbar} (HA(t) - A(t)H) = \{H, A(t)\}_h.$$

Introducendo l'operatore $U(t) = e^{-\frac{i}{\hbar}Ht}$ detto anche *operatore di evoluzione*, osservato che

$$U^* = \left(e^{-\frac{i}{\hbar}Ht}\right)^* = e^{\left(-\frac{i}{\hbar}Ht\right)^*} = e^{\frac{i}{\hbar}Ht}$$

la soluzione (24) si riscrive nella forma

$$A(t) = U^*AU.$$

Per risalire alla formulazione di Liouville della meccanica valutiamo

$$\begin{aligned}(A(t)|w) &= Tr(MA(t)) = Tr(MU^*AU) = Tr(U^*AUM) = \\ &= Tr(AUMU^*) = Tr(AM(t)) = (A|M(t))\end{aligned}$$

dove abbiamo definito l'operatore densità variabile

$$M(t) := UMU^*.$$

L'equazione di evoluzione corrispondente assume la forma

$$\begin{aligned}M'(t) &= U'MU^* + UM(U^*)' = -\frac{i}{\hbar}HUMU^* + \frac{i}{\hbar}UMU^*H = \\ &= -\frac{i}{\hbar}(HM - MH) = -\{H, M\}_h.\end{aligned}\tag{25}$$

Nel caso di uno stato puro $A = P_\psi$ la soluzione (24) assume la forma

$$P_\psi(t) = UP_\psi U^*.$$

In meccanica quantistica si è soliti ricercare la dipendenza rispetto al tempo nella sola evoluzione dello stato, ovvero facendo in modo che sia $P_\psi(t) = P_{\psi(t)}$. Una scelta comune per ottenere ciò è considerare l'evoluzione lineare $\psi(t) = U\psi$. Infatti, con tale scelta, per il generico vettore φ risulta

$$P_{\psi(t)}\varphi = \langle\psi(t), \varphi\rangle\psi(t) = \langle U\psi, \varphi\rangle U\psi = U(\langle\psi, U^*\varphi\rangle\psi) = U(P_\psi(U^*\varphi)) = P_\psi(t)\varphi.$$

Con queste scelte perveniamo all'usuale espressione dell'equazione di Schrodinger

$$\frac{d\psi(t)}{dt} = U'\psi = -\frac{i}{\hbar}HU\psi = -\frac{i}{\hbar}H\psi(t) \Rightarrow i\hbar\frac{d\psi(t)}{dt} = H\psi.$$

Bibliografia

- [1] L. D. Faddeev, O. A. Yakubovskii, *Lectures on Quantum Mechanics for Mathematics Students*, AMS, 2009.
- [2] L. Granieri, *Ottimo in Matematica*, La Dotta, 2016.
- [3] L. Granieri, *Sul Teorema Fondamentale dell'Algebra*, Periodico di Matematiche, Mathesis, N.1 2019.
- [4] L. Granieri, *Sostituisci e parti*, Archimede N. 4, 2018.
- [5] J. C. Polkinghorne, *Il Mondo dei Quanti*, Garzanti, 1986.

Crittografia: una Prospettiva Didattica – parte 2

Cryptography: an Educational Perspective – part 2

Nicola Fusco¹

Abstract

Cryptography is the ancient discipline that deals with developing methods for transforming text messages so that they are incomprehensible to anyone other than the legitimate recipient. In this work we will look at a small selection of modern cryptographic techniques that have links with high school mathematics topics. Given the importance that cryptography has assumed in the ubiquitous digital communications, this topic can constitute a module of civics education.

Introduzione

Ogni tecnica crittografica consiste nell'operazione di *cifratura*, con cui un messaggio in chiaro viene trasformato in un messaggio incomprensibile, e in quella inversa di *decifratura*. La *crittografia a chiave pubblica* raccoglie tutte tecniche sviluppate dal 1977, caratterizzate da un uso esteso di algoritmi coinvolgenti grandi numeri interi, la cui gestione rende necessario l'uso di computer. Nella crittografia a chiave pubblica le chiavi di codifica e decodifica sono diverse (da ciò il nome alternativo di *crittografia asimmetrica*) per cui la prima è resa pubblica e la seconda è mantenuta segreta a tutti, anche al mittente del messaggio.² Nel seguito sono esposti un metodo di cifratura a chiave pubblica (il cifrario RSA) e un metodo (la Zero Knowledge Proof) che permette di controllare il possesso di una certa informazione da parte di un soggetto, senza che tale informazione debba viaggiare tra i due interlocutori. In tutti i casi sono evidenziati i collegamenti con argomenti di matematica svolti nelle scuole secondarie di II grado. Le appendici sono disponibili presso l'autore.

¹Docente di Matematica e Fisica presso il Liceo Scientifico Statale "A. Scacchi" Bari (BA), fusco.nicola@liceoscacchibari.it.

²Tali metodi crittografici si sono rivelati indispensabili in quanto la stessa esistenza dei computer, che li ha resi possibili, ha reso estremamente vulnerabile tutta la crittografia a chiave privata che al giorno d'oggi non è più utilizzata da sola. Continua ad essere usata solo come completamento di tecniche a chiave pubblica.

1. Cifrario RSA

Il cifrario RSA³ si basa sull'*aritmetica modulare*⁴.

Bob deve inviare un messaggio riservato ad Alice. I due allora eseguono i passaggi schematizzati nella seguente tabella, in cui si distingue tra le operazioni il cui risultato è tenuto nascosto da quelle in cui il risultato è invece pubblicato (in termini generali o perché inviato all'altra persona e quindi suscettibile di intercettazione): la prima colonna indica chi esegue le operazioni indicate nella riga, la seconda e terza colonna mostrano, rispettivamente, le operazioni private e pubbliche.⁵

<i>Operatore</i>	<i>Operazione Privata</i>	<i>Operazione Pubblica</i>
Alice	Sceglie due primi, p e q . Es. $p = 5, q = 11$	Calcola il prodotto tra i due primi $N = p \cdot q$, Es. $N = 5 \cdot 11 = 55$
Alice	Calcola la funzione di EULERO di N $b = \varphi(N) = (p - 1)(q - 1)$ Es. $b = 4 \cdot 10 = 40$	Determina la chiave di cifratura: il minimo intero coprimo con b $e = \min\{x \in \mathbb{N}: MCD(b, x) = 1\}$ Es. $e = 3$
Alice	Calcola la chiave di decifratura: l'inverso di e modulo b $d: [d \cdot e]_b = 1$, Es. $d = 27$	
Bob	Converte il messaggio da inviare in una sequenza di numeri m_i minori di N . Es. $m = 7$	Cifra m nel numero c usando e e N e lo invia ad Alice $c = [m^e]_N$ Es. $c = [7^3]_{55} = 13$
Alice	Decifra il numero c nel numero m usando d e N $m = [c^d]_N$ Es. $m = [13^{27}]_{55} = 7$	

Tabella 4: Schema delle operazioni del Cifrario RSA.

L'operazione di decifratura funziona grazie al fatto che e e d sono inversi tra loro nell'aritmetica modulo b e nell'aritmetica modulo n gli esponenti si riducono modulo $\varphi(n)$. Quindi $[c^d]_N = [(m^e)^d]_N = [m^{ed}]_N = [m^1]_N = m$.

³Dalle iniziali dei cognomi dei suoi inventori: Rivest, Shamir e Adleman.

⁴Gli elementi fondamentali di aritmetica modulare sono riportati in appendice.

⁵Nel seguito, per maggiore sintesi, si usa la notazione $[a]_n$ per indicare $a \bmod n$.

1.1 Vantaggi e svantaggi di questo metodo

Nelle applicazioni reali i numeri p e q hanno un altissimo numero di cifre, perché l'unico modo per violare il cifrario RSA è scomporre N , operazione che tuttavia richiede un tempo che cresce esponenzialmente con il numero di cifre dei numeri primi coinvolti. Pertanto la pubblicazione di N , purché esso sia ottenuto da numeri primi molto grandi, non crea problemi perché nessuno è in grado di scomporlo nel tempo in cui è necessario mantenere la segretezza del messaggio. Il crescere della potenza di calcolo dei computer permette di accorciare il tempo necessario alla scomposizione, ma questo problema può essere aggirato senza difficoltà aumentando la grandezza dei numeri primi utilizzati:⁶ si può dimostrare che il tempo necessario alla scomposizione cresce in modo esponenziale con il numero di cifre dei numeri coinvolti. Questo rende il cifrario RSA di fatto inviolabile.

Lo svantaggio è che richiede un numero elevato di calcoli per ogni cifratura e di conseguenza, se lo si volesse utilizzare per cifrare direttamente un messaggio, le comunicazioni risulterebbero oltremodo rallentate. Per tale motivo il cifrario RSA viene solitamente usato o per comunicazioni molto brevi (ad esempio per trasmettere un codice di autorizzazione di una operazione bancaria) o solo per trasmettere una chiave privata e poi il messaggio vero e proprio viene cifrato con un sistema a chiave privata. In questo modo il tempo necessario si abbatte tantissimo e ciò rende possibile cambiare chiave privata con una frequenza molto alta, anche ogni pochi minuti, trasferendo all'intera comunicazione l'invulnerabilità del cifrario RSA.

1.2 Argomenti di matematica connessi

Questo sistema di cifratura è collegato direttamente con tutta l'aritmetica dei numeri naturali: numeri primi, divisori, scomposizione in fattori primi, divisione con resto, proprietà delle potenze.

2. Zero Knowledge Proof

Le ZKP sono una delle ultime grandi innovazioni in termini di comunicazione sicura benché ormai il concetto risalga al 1985. Esse consistono in protocolli di scambi di domande e risposte, volte a permettere ad un soggetto verificatore (Viola nel seguito) di assicurarsi che un altro soggetto dimostratore (Daniela nel seguito) sia realmente in possesso di una data informazione (ad esempio la password per accedere ad un account), senza che il dimostratore trasmetta effettivamente tale informazione.

Non si tratta quindi di un vero e proprio metodo crittografico, perché nel processo non viene trasmesso alcun messaggio nascosto da Viola a Daniela, ma in molti casi può essere sufficiente accertarsi che una data affermazione sia vera con una certa affidabilità e risparmiarsi la trasmissione di un'informazione che potrebbe fornire la certezza della veridicità ma nello stesso tempo esporre tale informazione al rischio di intercettazione.

⁶Questo spiega l'impegno che viene profuso nel ricercare numeri primi sempre più grandi: non è una mera curiosità matematica, ma una necessità tecnologica per mantenere la sicurezza delle comunicazioni.

In generale una ZKP deve soddisfare tre requisiti:

Completezza: se l'affermazione è vera Daniela deve essere sempre in grado di dimostrare tale verità a Viola.

Correttezza: se l'affermazione è falsa Viola deve essere potenzialmente in grado di scoprirne la falsità.

Conoscenza zero: Viola è solo in grado di acquisire informazioni sulla verità o falsità dell'affermazione ma non sul contenuto dell'affermazione.

Di seguito sono descritti tre esempi che mostrano come sia possibile trasferire il valore di verità di una data affermazione senza inviare alcunché dell'affermazione in sé. Il terzo esempio è proprio un'applicazione realistica di questi protocolli. Tuttavia neanche il terzo di questi esempi descrive un reale protocollo ZKP: quelli realmente usati nell'ambito delle comunicazioni digitali sono alquanto complessi e la loro descrizione andrebbe ben al di là degli obiettivi di questo articolo.

Esempio 1: le sfere colorate

Un esempio classico con cui si mostra la possibilità di realizzare una cosa del genere è la seguente: immaginiamo che Daniela dia a Viola due sfere identiche in tutto tranne che per il colore: una è rossa e una è verde. Daniela dice a Viola cosa gli sta dando, ma Viola non è in grado di controllarlo in autonomia perché è daltonica, quindi non è in grado di distinguere il rosso dal verde e pertanto non sa se Daniela gli stia dicendo la verità o in realtà le due sfere siano di colore uguale. Cosa può fare Viola per assicurarsi con un certo grado di sicurezza che Daniela non le stia mentendo?

Viola procede nel modo seguente: nasconde le sfere dietro la schiena, dicendo a Daniela che potrebbe stare scambiando le sfere di mano oppure no, poi le rende nuovamente visibili e chiede a Daniela di dire se ha effettivamente scambiato le sfere. E ripete questo processo più volte, diciamo n .

Qualunque sia il colore delle sfere, Viola sa con certezza se le ha scambiate oppure no. Se le due sfere hanno effettivamente colore diverso, anche Daniela può sapere con certezza se è avvenuto lo scambio oppure no, perché ad ogni processo le basta vedere se ogni mano di Viola regge la sfera dello stesso colore che aveva prima oppure no. Se le due sfere hanno invece lo stesso colore Daniela non sa se lo scambio è avvenuto o no, e può solo tirare a indovinare, con probabilità $\frac{1}{2}$ di indovinare o di sbagliare ogni volta.

Supponiamo che Viola effettui n volte il processo e ponga ogni volta la domanda a Daniela. Se Daniela sbaglia almeno una volta, non appena commette l'errore Viola sa con certezza che Daniela stava mentendo perché, come detto, se le sfere sono diverse Daniela non può sbagliare in quanto vede senza problemi che colore sia tenuto in ciascuna mano da Viola. Ma se Daniela risponde correttamente tutte le volte, Viola ha quindi la certezza che Daniela stia dicendo la verità? No, perché Daniela potrebbe aver indovinato per puro caso tutte le varie volte e quindi aver mentito senza che Viola se ne sia potuta accorgere. Tuttavia Viola può stimare il grado di affidabilità dell'affermazione iniziale di Daniela usando il teorema di BAYES.

Supponiamo che Viola a priori non abbia alcun motivo né per dubitare né per fidarsi di Daniela, quindi attribuisce inizialmente uguale probabilità sia all'evento F = "Daniela ha mentito" sia all'evento V = "Daniela è sincera": $P(F) = P(V) = \frac{1}{2}$.

Se Daniela è sincera la probabilità che indovini tutte le n risposte (evento I) è 1, $P_n(I|V) = 1$, mentre che sbagli anche solo una volta (evento E) è 0, $P_n(E|V) = 0$. Se invece Daniela mente la probabilità che indovini tutte le n risposte (evento I) è una probabilità congiunta di n eventi indipendenti ciascuno con probabilità $\frac{1}{2}$ di realizzarsi, $P_n(I|F) = (\frac{1}{2})^n$, mentre che sbagli anche solo una volta (evento E) è la probabilità complementare, $P_n(E|F) = 1 - (\frac{1}{2})^n$.

Con il teorema di BAYES Viola può quindi determinare le probabilità a posteriori che Daniela sia sincera o menta in base a se abbia risposto correttamente alle n domande o ne abbia sbagliata almeno una.

$$P_n(V|I) = \frac{P_n(I|V) P(V)}{P_n(I|V) P(V) + P_n(I|F) P(F)} = \frac{1}{1 + 2^{-n}}$$

$$P_n(F|I) = 1 - P_n(V|I) = \frac{1}{2^n + 1}$$

$$P_n(V|E) = \frac{P_n(E|V) P(V)}{P_n(E|V) P(V) + P_n(E|F) P(F)} = 0$$

$$P_n(F|E) = 1 - P_n(V|E) = 1$$

Come vediamo, quindi, Viola è in grado di stabilire con certezza che Daniela mente se sbaglia almeno una volta, o di attribuire una fiducia crescente⁷ man mano che Daniela risponde correttamente ad un numero crescente di domande. Tutto questo avviene benché Viola resti sempre ignara del colore delle sfere che ha in mano. Quindi a Viola non viene trasmessa nessuna informazione aggiuntiva se non la veridicità o falsità dell'affermazione iniziale di Daniela.

Esempio 2: il sudoku

Daniela afferma di aver risolto il sudoku di un concorso di un giornale e vuole vendere la soluzione a Viola. Naturalmente Viola vuole la prova che Daniela possieda davvero la soluzione prima di pagare, ma altrettanto naturalmente Daniela vuole essere pagata prima di mostrare la soluzione. Daniela ha quindi bisogno di fornire una prova a conoscenza zero del fatto che possiede realmente la soluzione.

Viola disegna su un grande foglio una griglia da sudoku di 81 caselle abbastanza grandi da contenere una carta da gioco e poi Viola piazza, a faccia in su in ciascuna delle caselle contenenti i numeri noti a priori, tre carte di valore uguale e corrispondente al numero da 1 a 9 che il giornale ha stampato nella traccia del sudoku in quella casella.

⁷Nell'espressione $\frac{1}{1+2^{-n}}$, che fornisce $P_n(V|I)$, al crescere di n l'esponenziale 2^{-n} diminuisce sempre di più, fino a diventare piccolissimo per valori di n abbastanza grandi. Di conseguenza l'intera espressione $\frac{1}{1+2^{-n}}$ aumenta a partire dal valore $\frac{1}{2}$ avvicinandosi sempre di più al valore 1. Del resto $P_n(F|I)$, per motivi analoghi, diminuisce all'aumentare di n dal valore $\frac{1}{2}$ avvicinandosi sempre di più al valore 0, dato che l'addendo al denominatore 2^n aumenta all'aumentare di n fino a diventare enorme per valori di n abbastanza grandi.

Daniela successivamente piazza su ciascuna casella vuota (corrispondente ad una casella bianca nella griglia iniziale del sudoku non risolto) tre carte di uguale valore corrispondente al numero che lei ha scritto in tale casella nella sua soluzione. Queste carte vengono piazzate a faccia in giù, in modo che nessuno ne conosca il valore a parte Daniela.

Dopodiché Viola prende a caso una carta da ogni casella della prima riga e le mette in una busta, che viene chiusa. Poi fa lo stesso con tutte le altre righe, con tutte le colonne e con tutti i quadranti. Nel fare questo capovolge le carte che aveva piazzato a faccia in su, in modo che nella busta non si possa capire quale fossero le carte dei numeri noti a priori e quelle dei numeri della soluzione. Alla fine del processo ci sono quindi 27 buste ciascuna contenente 9 carte. Si mescolano tutte le buste e le si apre per esaminarne il contenuto.

Una soluzione corretta del sudoku genera necessariamente, tramite la procedura descritta, 27 buste contenenti ciascuna 9 carte numerate da 1 a 9. Se Viola trova anche una sola delle 27 buste che non contiene tutti i valori da 1 a 9 allora sa con certezza che Daniela non possiede realmente una soluzione del sudoku proposto. Se invece trova che tutte le 27 buste hanno un contenuto corretto allora può avere ragionevole fiducia che Daniela davvero conosca la soluzione. Naturalmente in questo secondo caso non si può avere la certezza perché Daniela potrebbe avere prodotto le 27 buste col contenuto corretto anche per caso. Ad esempio Daniela potrebbe non aver messo carte con lo stesso valore in ogni casella, ma aver messo tre carte diverse in modo che una si "sistemasse" bene rispetto alla riga di appartenenza, una rispetto alla colonna e una rispetto al quadrante (ma, in generale, non tutte andassero bene per tutte e tre le richieste) e Viola, pur scegliendo a caso, ha preso di volta in volta proprio la carta che si allinea bene con tutte le altre della stessa riga, colonna o quadrante.

Qual è la probabilità che le 27 buste siano corrette per caso?

Se Daniela non ha risolto il sudoku e ha messo in ogni casella 3 carte come descritto poco sopra, allora c'è una probabilità pari a $\frac{1}{6}$ che Viola per caso le peschi nel modo "giusto" da una singola casella.⁸ Le caselle in totale sono 81 e indichiamo con C il numero di caselle dal contenuto inizialmente noto. Quindi la probabilità che Viola peschi per caso tutte le carte nell'ordine giusto è $(\frac{1}{6})^{81-C}$ e questa coincide con la probabilità che la prova singola sia favorevole a Daniela anche se essa ha mentito. Quindi, usando la stessa notazione dell'esempio precedente, abbiamo

$$P_C(I|V) = 1 \wedge P_C(E|V) = 0$$

$$P_C(I|F) = \left(\frac{1}{6}\right)^{81-C} \wedge P_C(E|F) = 1 - \left(\frac{1}{6}\right)^{81-C}$$

da cui

$$P_C(V|I) = \frac{1}{1 + 6^{-81+C}} \wedge P_C(F|I) = \frac{1}{6^{81-C} + 1}$$

$$P_C(V|E) = 0 \wedge P_C(F|E) = 1$$

⁸3 carte si possono ordinare in $3! = 6$ modi diversi, ma solo uno di questi ordinamenti è quello giusto.

Se a Viola non basta un singolo tentativo può far ripetere la procedura n volte. In tal caso le probabilità condizionate iniziali sono

$$P_{C,n}(I|V) = 1 \wedge P_{C,n}(E|V) = 0$$

$$P_{C,n}(I|F) = \left(\frac{1}{6}\right)^{n(81-C)} \wedge P_{C,n}(E|F) = 1 - \left(\frac{1}{6}\right)^{n(81-C)}$$

mentre quelle ottenute dal teorema di BAYES sono

$$P_{C,n}(V|I) = \frac{1}{1 + 6^{-n(81-C)}} \wedge P_{C,n}(F|I) = \frac{1}{6^{n(81-C)} + 1}$$

$$P_{C,n}(V|E) = 0 \wedge P_{C,n}(F|E) = 1$$

Anche in questo caso Viola è in grado di stabilire con certezza che Daniela mente se sbaglia almeno una volta, o di attribuire una fiducia crescente man mano che Daniela dispone correttamente un numero crescente di schemi risolutivi. E tutto questo avviene benché Viola resti sempre ignara di quale sia la soluzione, perché dal contenuto delle buste, essendo state mescolate, non si può risalire allo schema che le ha prodotte.

Esempio 3: la password

Daniela è l'app con cui accedere all'account di una data banca, l'account è protetto da una password di 10 caratteri. I caratteri possibili sono tutti quelli disponibili sulla tastiera (con la distinzione tra maiuscole e minuscole), cioè 102 possibili simboli diversi. Viola è il software di controllo delle password.

Naturalmente l'invio della password, anche se criptata, espone comunque al rischio di intercettazione e quindi di decifrazione, per cui Daniela e Viola sono state programmate per attuare il controllo mediante una ZKP. Daniela chiede all'utente di inserire la password ma non la invia in alcun modo, memorizza solo, per ciascun carattere, un valore numerico (da 1 a 102, oppure da 99 a 200 se si vuole evitare di usare numeri troppo piccoli) che è stato attribuito in precedenza ad ogni carattere. Viola invia a Daniela una sequenza ordinata di 4 numeri (n_1, n_2, n_3, n_4) da 1 a 10 (ciascuno dei quali indica una posizione di un carattere della password) e Daniela deve inviare in risposta il risultato P del seguente calcolo modulare

$$P = \left[(N_1 - N_2)^{N_3} \right]_{N_4}$$

dove N_1 è il valore numerico del carattere che si trova in posizione n_1 , N_2 è il valore numerico del carattere che si trova in posizione n_2 , e così via. Questa richiesta può essere svolta numerose volte, perché il numero di disposizioni di 4 caratteri che si possono formare con 10 caratteri è $\frac{10!}{(10-4)!} = 5040$.

Dal risultato del calcolo non è possibile risalire ai quattro numeri, neanche se si conoscessero i valori P calcolati con gli stessi numeri ma usati in svariati ordini diversi. Per cui anche se i valori P inviati da Daniela a Viola vengono intercettati, questo non permette la ricostruzione della password, neanche dopo numerose intercettazioni.

Anche in questo caso un solo errore implica il riconoscimento della falsità: chi possiede realmente la password non può sbagliare nessuno dei calcoli richiesti. Nello stesso tempo ogni risposta andata a buon fine aumenta la fiducia che Daniela possieda realmente la password dell'account. La probabilità di indovinare il risultato di una singola operazione è (nell'ipotesi di usare valori dei caratteri da 99 a 200) $\frac{1}{200}$: chi non possiede la password non conosce neanche il numero N_4 rispetto a cui viene fatta la riduzione modulare⁹ e pertanto, dal suo punto di vista, il valore massimo che può avere P è in ogni caso il valore massimo che può assumere la base di riduzione modulare. Quindi le probabilità condizionate iniziali, dopo n domande, sono

$$P_n(I|V) = 1 \wedge P_n(E|V) = 0 \wedge P_n(I|F) = \left(\frac{1}{200}\right)^n$$

$$P_n(E|F) = 1 - \left(\frac{1}{200}\right)^n$$

mentre quelle ottenute dal teorema di BAYES sono

$$P_{C,n}(V|I) = \frac{1}{1 + 200^{-n}} \wedge P_{C,n}(F|I) = \frac{1}{200^n + 1}$$

$$P_{C,n}(V|E) = 0 \wedge P_{C,n}(F|E) = 1$$

Anche in questo caso Viola è in grado di stabilire con certezza che Daniela mente se sbaglia almeno una volta, o di attribuire una fiducia crescente man mano che Daniela invia correttamente un numero crescente di risultati di calcoli modulari. L'unica differenza rispetto agli esempi precedenti è che in questo caso Viola conosce l'informazione di cui deve verificare il possesso da parte di Daniela ma tale informazione non viene mai trasmessa da Daniela a Viola e pertanto non è suscettibile di intercettazione.

2.1. Vantaggi e svantaggi di questo metodo

Il vantaggio principale, come già detto, è che non c'è alcuna trasmissione di informazioni utili per un intercettatore fraudolento, quindi un qualunque protocollo può essere riutilizzato a piacere senza che rischi di violazione.

Il secondo vantaggio è che si ha una notevole sicurezza nelle trasmissioni pur non usando alcun protocollo crittografico, e questo naturalmente velocizza le trasmissioni grazie all'aggiramento delle procedure di codifica e decodifica.

Gli svantaggi sono che i protocolli effettivamente usati richiedono un alto numero di risorse (capacità di calcolo e memoria) rispetto ai sistemi crittografici standard, per cui sono implementabili solo su dispositivi abbastanza potenti.¹⁰

⁹Questo numero, peraltro, può essere aumentato a piacere dato che c'è completa libertà nella composizione della tabella di corrispondenza tra i 102 caratteri e i valori numerici usati. Non è neanche necessario che i numeri siano consecutivi.

¹⁰Ad esempio, gli smartphone attualmente in commercio non sono in grado di gestire questi protocolli in tempi ragionevolmente brevi per un loro utilizzo.

Inoltre non sono sistemi che permettono lo scambio di informazione in modo sicuro, per cui sono utilissimi per controllare, ad esempio, il possesso da parte di un soggetto delle giuste credenziali di accesso ad un dato account, ma non possono poi essere usati per cifrare le comunicazioni che il soggetto deve inviare al sistema gestore dell'account successivamente.

Infine, per permettere il raggiungimento di alti valori di affidabilità nella verifica del possesso di informazioni complesse, questi protocolli richiedono la modellizzazione di strutture matematiche altamente complesse, la cui conoscenza, unita alla competenza informatica necessaria all'implementazione, appartiene al know-how di poche persone al mondo, che sostanzialmente dominano questo settore della sicurezza. Naturalmente ciò comporta che questo settore è una delle nicchie più promettenti dal punto di vista lavorativo e di sviluppo imprenditoriale nell'ambito dell'alta tecnologia.

2.2. Argomenti di matematica connessi

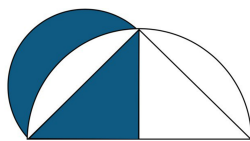
Le ZKP sono evidentemente collegate con tutta la teoria elementare della Probabilità, in particolare la formula di Leibniz per gli eventi equiprobabili, la probabilità congiunta, in particolare di eventi indipendenti, la probabilità condizionata e il teorema di BAYES.

Questo argomento ha anche una connessione con il metodo scientifico, soprattutto nel suo sviluppo popperiano. Esattamente come nel metodo scientifico, in questo caso si procede per prove, confutazioni e corroborazioni: il fallimento di una prova comporta la scoperta della falsità dell'affermazione iniziale, esattamente come per una teoria scientifica, che è trovata fallace non appena fallisce una prova sperimentale, mentre gli esiti positivi delle prove non forniscono mai una conferma definitiva, aumentano solo la fiducia nella veridicità dell'affermazione, così come le lunghe sequenze di esperimenti in accordo con una teoria la corroborano senza mai verificarla o confermarla in modo assoluto.¹¹

Bibliografia

- [1] BONA VOGLIA P., La Crittografia da Atbash a RSA, 2022, www.crittologia.eu.
- [2] GOLDREICH O., KRAWCZYK H., On the Composition of Zero-Knowledge Proof Systems, ICALP 1990. Lecture Notes in Computer Science, vol 443, pp 268-282, Springer, doi.org/10.1007/BFb0032038.
- [3] GRADWOHL R., NAOR M., PINKAS B., ROTHBLUM G., Cryptographic and Physical Zero-Knowledge Proof Systems for Solutions of Sudoku Puzzles, www.wisdom.weizmann.ac.il/naor/PAPERS/sudoku_abs.html.

¹¹Per approfondire questo aspetto si può leggere Fusco N., Probabilità ed Epistemologia di Popper, Periodico di matematiche vol. 10 serie XIV anno CXXVIII, n. 2 Mag-Ago 2018, pp. 75-88, ISSN: 1582-8832.



Mathesis

Presidente

Domenica Di Sorbo

È presidente dal 24 febbraio 2024.

presidente-disorbo@mathesisnazionale.it

Consiglio Nazionale

Salvatore Cuomo *vice presidente*, Maria Coccozza *segretaria*, Roberto Franzina *tesoriere*, Serenella Iacino, Vincenzo Iorfida, Mina Lavopa, Susy Osti, Marcello Pedone, Francesco Sicolo, Pasqualina Ventrone *consiglieri*.

Sezioni

Abruzzo – Avellino – Bari – Brescia – Caserta – Castellamare di Stabia – Catania – Crotone – Ferrara – Firenze – Grottaglie – Latina – Lecce – Mantova – Messina – Milano – Mondragone – Napoli – Napoli Flegrea – Olbia – Roma – Rovigo – Salerno – Serra San Bruno – Spoleto – Terni – Udine – Varese – Venezia – Verona – Vicenza.

Rivista

Periodico di Matematiche

Dal 1886 Periodico di Matematica, dal 1921 "di Matematiche". La proprietà della testata è della Mathesis dal 2012 per concessione dell'Editore Zanichelli.

Sito web

www.mathesisnazionale.it



In memoria di **Francesco de Giovanni**,
già Direttore del Periodico di Matematiche,
Presidente della Mathesis Nazionale,
Professore Ordinario di Algebra
all'Università degli Studi di Napoli Federico II.

MATHESIS

Società Italiana di Scienze Matematiche e Fisiche

c/o Dipartimento di Matematica e Fisica
Università degli Studi della Campania "Luigi Vanvitelli"
Viale Abramo Lincoln 5 - 81100 Caserta (CE)
www.mathesisnazionale.it • info@mathesisnazionale.it