

Numero 10 - 1996

RATIO MATHEMATICA

**Rivista di
applicazioni della matematica**

a cura di

Franco Eugeni e Mario Gionfriddo

Comitato Scientifico

Albrecht Beutelspacher, Giessen

Piergiulio Corsini, Udine

Bal K. Dass, Delhi

Franco Eugeni, Roma

Mario Gionfriddo, Catania

Antonio Maturo, Pescara

Nicola Rotondale, L'Aquila

Aniello Russo Spina, L'Aquila

Romano Scozzafava, Roma

Maria Tallini Scafati, Roma

Numero 10 - 1996

RATIO MATHEMATICA

Rivista di
applicazioni della matematica

a cura di

Franco Eugeni e Mario Gionfriddo

Comitato Scientifico

Albrecht Beutelspacher, *Giessen*

Piergiulio Corsini, *Udine*

Bal K. Dass, *Delhi*

Franco Eugeni, *Roma*

Mario Gionfriddo, *Catania*

Antonio Maturo, *Pescara*

Nicola Rotondale, *L'Aquila*

Aniello Russo Spena, *L'Aquila*

Romano Scozzafava, *Roma*

Maria Tallini Scafati, *Roma*

Numero 10 - 1996

RATIO MATHEMATICA

Rivista di
applicazioni della matematica

INDICE

F. EUGENI, F. ZUANNI, <i>Un problema di minimo per il nastro di Möbius</i>	pag.	3
S. FERRI, <i>Elementi di geometrie finite</i>	pag.	9
M. ZANNETTI, <i>Il vantaggio di operare in un ambiente finito</i> ..	pag.	23
C. D'ANIELLO, <i>Bruno Rizzi e il Fusionismo</i>	pag.	31
C. DI FOGGIA, R. PROSPERI, <i>Rappresentazione vettoriale delle sinusoidi</i>	pag.	35
F. DI GENNARO, <i>Campi di Galois: una presentazione divulgativa</i>	pag.	43
F. MANCINELLI, <i>Note su due problemi di geometria e successioni ricorsive</i>	pag.	53
F. PASCUCCI, <i>Aspetti combinatori della logica</i>	pag.	61
A. MATURO, B. FERRI, <i>Un modello matematico per la valutazione di fattibilità dei programmi integrati d'intervento urbano</i>	pag.	67
L. CERRITELLI, F. MERCANTI, <i>Algebra astratta e insegnamento</i>	pag.	85
F. MERCANTI, L. CERRITELLI, V. DI MARCELLO <i>Un nuovo approccio alla teoria dei Sistemi Digitali Multicanale (SDM)</i>	pag.	91
E. GELSOMINI, <i>Un ipergruppo associato all'insieme delle matrici quadrate di ordine n non singolari</i>	pag.	99

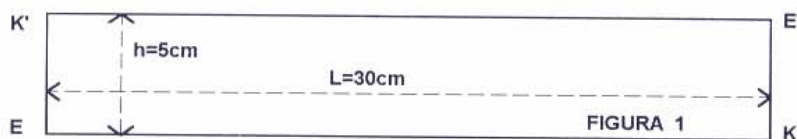
UN PROBLEMA DI MINIMO PER IL NASTRO DI MÖBIUS

Franco Eugeni*, Fulvio Zuanni**

SUNTO - Un nastro di Möbius si costruisce a partire da una striscia di carta rettangolare di lato "corto" h e lato "lungo" L . Fissato h si pone il problema di trovare il minimo di L . Sperimentalmente si può intuire che tale valore minimo è $L_{\min} = h\sqrt{3}$, ma il dimostrarlo non è banale ed accosta l'allievo a problematiche che presentano i tre aspetti di analisi sperimentale, formulazione del modello e risoluzione logico-deduttiva.

UN MODELLO MATERIALE DI NASTRO DI MÖBIUS

Il nastro di Möbius è un ben noto esempio di superficie ad una sola faccia. La costruzione classica del nastro di Möbius è ben nota. Prendiamo una striscia di carta rettangolare $EKE'K'$ di lati $h=5\text{cm}$ e $L=30\text{cm}$.



Si può piegare la carta in modo da formare un cilindro facendo coincidere E con K e K' con E' . Se invece si congiungono le due estremità dopo averne rovesciata una, facendo coincidere E con E' e K con K' , si ottiene un nastro di Möbius. Sperimentalmente, si può ripetere la costruzione accorciando L . Ciò si può fare agevolmente fino a circa $L=3h$. Con molta difficoltà si riesce a farlo ancora quando $L=2h$, mentre sembra impossibile con $L=h$.

* Facoltà di Ingegneria - Terza Università di Roma

** Facoltà di Ingegneria - Università de L'Aquila

Ricordiamo un modo semplice di costruire il modello materiale di nastro di Möbius che ci sarà utile per la risoluzione del problema. Prendiamo ancora una striscia di carta di altezza data h e lunghezza L "sufficientemente" grande. Pieghiamo ($\alpha_1 > 0$) la striscia lungo il segmento AB come in figura 2. E' immediato convincersi che non si perde di generalità imponendo che sia $\alpha_1 < \pi/2$.

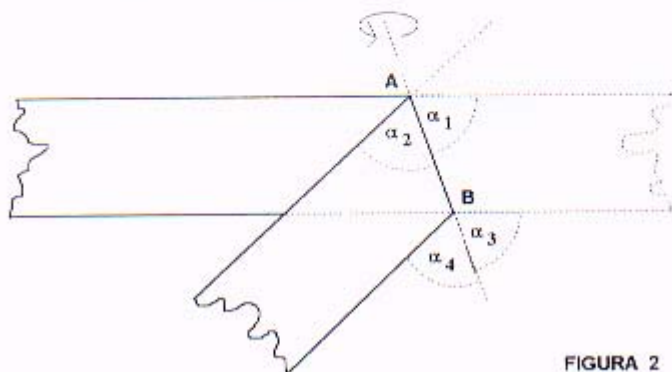


FIGURA 2

Gli angoli α_1 e α_2 sono tra loro congruenti per effetto del "piegamento". Inoltre gli angoli α_1 e α_2 sono congruenti rispettivamente agli angoli α_3 e α_4 perché corrispondenti.

Quindi, possiamo porre $\alpha := \alpha_1 = \alpha_2 = \alpha_3 = \alpha_4$ con $0 < \alpha < \pi/2$.

Ora, pieghiamo ($\beta_1 > 0$) la striscia lungo il segmento CD come in figura 3.

Gli angoli β_1 e β_2 sono tra loro congruenti per effetto del "piegamento".

Quindi, possiamo porre $\beta := \beta_1 = \beta_2$.

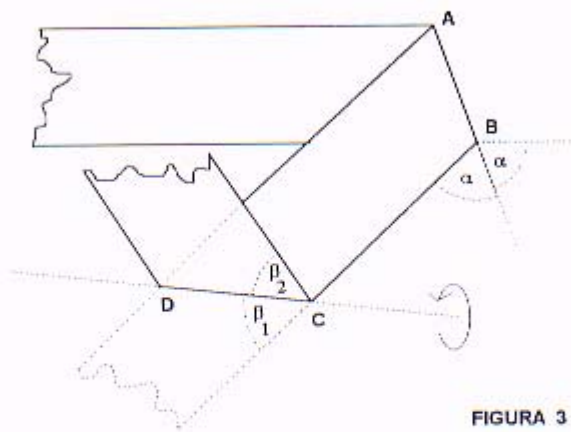


FIGURA 3

E' immediato rendersi conto che, affinché quest'ultimo pezzo di striscia si sovrapponga a quello iniziale, deve essere $\beta_1 < \pi/2$ e $(2\alpha + \beta_1 + \beta_2) = (2\alpha + 2\beta) > \pi$ da cui $(\pi/2 - \alpha) < \beta < \pi/2$.

A questo punto, tagliamo i due pezzi sovrapposti della striscia lungo il segmento EF, come in figura 4, ed uniamo i lembi così ottenuti.

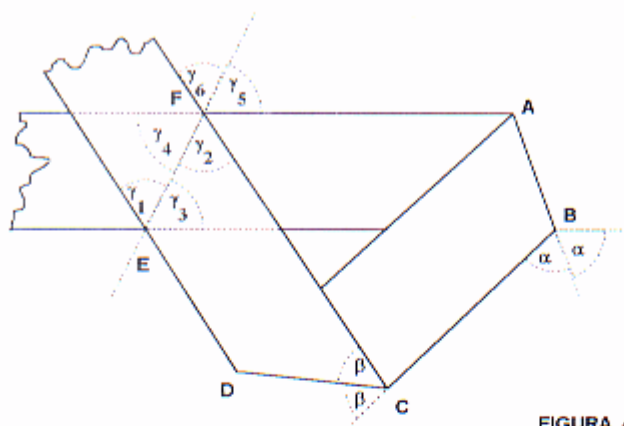


FIGURA 4

Con considerazioni analoghe alle precedenti si vede che $\gamma_1 = \gamma_2 = \gamma_3 = \gamma_4 = \gamma_5 = \gamma_6$.

Quindi, possiamo porre $\gamma := \gamma_1 = \gamma_2 = \gamma_3 = \gamma_4 = \gamma_5 = \gamma_6$.

Inoltre, si vede subito (figura 5) che $(2\alpha + 2\beta + 2\gamma) = 2\pi$, da cui

$$\gamma = (\pi - \alpha - \beta) \quad (1)$$

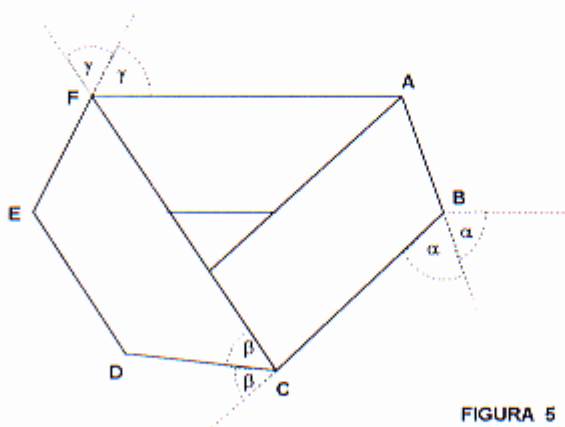


FIGURA 5

Il modello di nastro così costruito può essere "svolto" nello spazio tridimensionale "staccando" (letteralmente) i triangoli sovrapposti relativi ai lati AB, CD ed EF ottenendo la classica superficie ad una sola faccia che tutti conoscono. Se facciamo scorrere una penna, senza mai sollevarla, su tale oggetto riusciamo a tornare al punto di partenza e riusciamo a segnare tutto il nastro.

SULLA LUNGHEZZA DEL NASTRO DI MÖBIUS

Nel costruire il modello precedente, di altezza h data, abbiamo inizialmente richiesto che la striscia fosse "sufficientemente" lunga, cioè si è supposto $L \gg h$. In effetti, se utilizzassimo una striscia molto "corta", ad esempio con h ed L quasi uguali, ci accorgeremmo di non essere capaci di effettuare tale costruzione. Sorge spontaneo, quindi, il problema di stabilire quale debba essere, fissato h , il valore minimo di L , ovvero la lunghezza minima della striscia di carta, per poter ottenere l'indicato modello materiale di nastro di Möbius.

Un primo modo di "accorciare" la striscia è quello di far coincidere B con C (o, equivalentemente, D con E o A con F) come in figura 6.

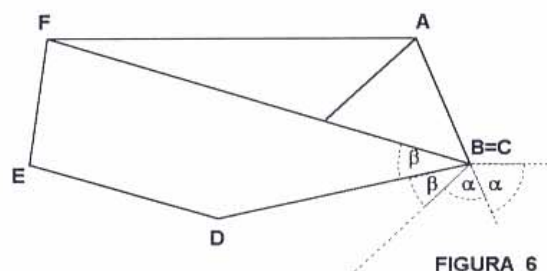


FIGURA 6

Questo è sempre possibile, mentre ci si accorge subito che, in generale, non si riesce a soddisfare a tutte e tre le condizioni contemporaneamente. Inoltre, si noti che deve essere:

$$2\alpha + \beta \leq 2\pi \quad (2)$$

$$\alpha + 2\beta \leq 2\pi \quad (3)$$

in caso contrario i punti E ed F appartenerebbero ad uno stesso lato della striscia iniziale e la costruzione descritta non sarebbe possibile.

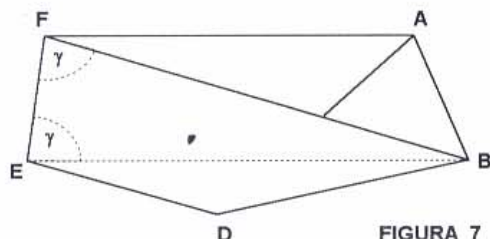
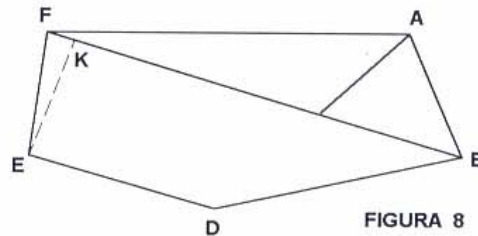


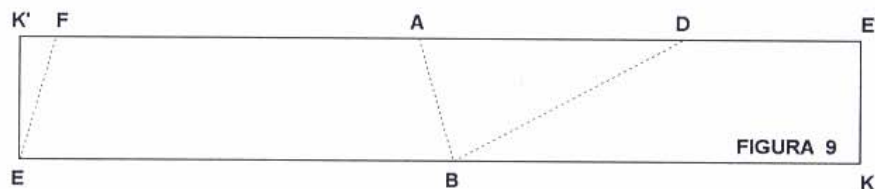
FIGURA 7

Dalla figura 7 si osserva che $BF=BE$ in quanto il triangolo EBF è isoscele sulla base EF .

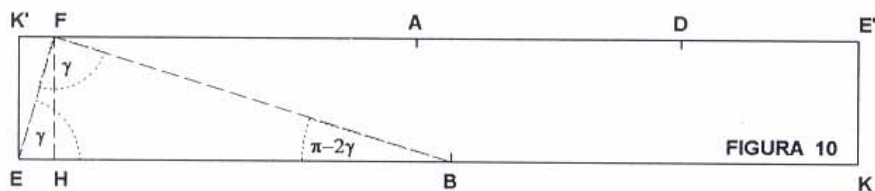
Indicata con K (figura 8) la proiezione di E su BF, abbiamo che $BK = BF - FK$.



Calcoliamo, ora, la lunghezza della striscia. Per farlo, tagliamo la striscia lungo il segmento EK. Svolgendola su di un piano otteniamo un rettangolo come in figura 9. **Si noti che** al segmento FK' della figura 9 corrisponde il segmento FK della figura 8, per cui $BK = BF - FK'$.



Indicata, come in figura 10, con H la proiezione di F su BE abbiamo, che $FK' = HE$ da cui $BH = BE - HE = BF - FK' = BK$.



Quindi, la lunghezza della striscia e' $L = HE + BH + BK = HE + 2BH$.

Dalla figura si vede subito che $FH = HE \tan \gamma$ e $FH = BH \tan(\pi - 2\gamma)$.

Da ben note formule di trigonometria si ha che

$$\tan(\pi - 2\gamma) = -\tan(2\gamma) = -\frac{2 \tan \gamma}{1 - \tan^2 \gamma} = \frac{2 \tan \gamma}{\tan^2 \gamma - 1}$$

$$\text{Per cui, si ha } HE = \frac{h}{\tan \gamma} \text{ e } 2BH = h \frac{\tan^2 \gamma - 1}{\tan \gamma} = h \tan \gamma - \frac{h}{\tan \gamma}.$$

Quindi, la lunghezza della striscia $EKE'K'$ (figure 9 e 10) e' data da

$$L = h \tan \gamma \quad (4)$$

UN PROBLEMA "NATURALE" DI MINIMO

Per minimizzare la lunghezza L appena trovata dobbiamo, evidentemente (poiché la tangente è crescente), minimizzare l'ampiezza dell'angolo γ . Tenendo conto di (1), (2) e (3) si ha

$$2\alpha + \beta \leq \pi \Rightarrow \alpha \leq \pi - \alpha - \beta \Rightarrow \alpha \leq \gamma$$

$$\alpha + 2\beta \leq \pi \Rightarrow \beta \leq \pi - \alpha - \beta \Rightarrow \beta \leq \gamma$$

$$\gamma = \pi - \alpha - \beta \Rightarrow 2\gamma = 2\pi - 2\alpha - 2\beta \Rightarrow 2\alpha + 2\beta - \pi = \pi - 2\gamma$$

$$2\alpha + \beta \leq \pi \Rightarrow 2\alpha + 2\beta - \pi \leq \beta \Rightarrow \pi - 2\gamma \leq \beta$$

$$\alpha + 2\beta \leq \pi \Rightarrow 2\alpha + 2\beta - \pi \leq \alpha \Rightarrow \pi - 2\gamma \leq \alpha$$

Quindi, possiamo scrivere

$$(\pi - 2\gamma) \leq \alpha \leq \gamma \quad (5)$$

$$(\pi - 2\gamma) \leq \beta \leq \gamma \quad (6)$$

Per (5) e (6) deve essere $(\pi - 2\gamma) \leq \gamma$, cioè $\gamma \geq \pi/3$. Quindi, la lunghezza minima si ottiene con $\gamma = \pi/3$ e, tenendo conto di (4), è

$$L_{\min} = h\sqrt{3}$$

In particolare, se $\gamma = \pi/3$ allora da (5) e (6) si ha $\alpha = \pi/3$ e $\beta = \pi/3$. Dunque il nastro di lunghezza minima si "ottiene" a partire da un triangolo equilatero.

Si noti che l'ultima relazione rappresenta più un estremo inferiore che un minimo per il modello materiale nel momento in cui si richiama "l'assioma intuitivo" di distacco dei triangoli sovrapposti altresì detto "possibilità" di svolgere il nastro nello spazio". In altre parole, il modello minimizzato "rimane schiacciato nel piano".

BIBLIOGRAFIA

1. CARL B. BOYER, *Storia della Matematica*, Oscar Mondadori, 1994.
2. H.B. GRIFFITHS, *Möbius Bands and the Didactics of joy* - a German Philosophy, *Mathematics in School* (1981), 14-18.
3. H. ZEITLER, *My story about the Moebius bands or How to discover mathematics in a joyful way*, Atti del Convegno Internazionale "Cultura Matematica e Insegnamento" nel decimo anniversario della scomparsa di Luigi Campedelli, Firenze 30 - 31 maggio e 1 giugno 1988, 49-67.

ELEMENTI DI GEOMETRIE FINITE

Stefania Ferri*

SUNTO - Si fa vedere come si può costruire la geometria analitica su un piano affine finito. Vengono riportati risultati classici che però sono stati esposti in modo da essere compresi dagli alunni delle Scuole Medie Secondarie.

INTRODUZIONE

La geometria analitica che si studia nelle Scuole Medie Superiori è la rappresentazione di punti, rette e curve nel piano affine reale. Cioè le coordinate dei punti, i coefficienti delle equazioni delle rette e delle curve sono numeri reali. Talvolta tale piano viene complessificato nel senso che si considerano anche i punti a coordinate numeri complessi (quando per esempio si vogliono trovare le intersezioni di una circonferenza con una retta esterna).

Si può però, come è noto, considerare il caso in cui tali coordinate e tali coefficienti appartengano, invece che al campo dei numeri reali, ad un campo finito, cioè costituito da un numero finito di elementi. In questa situazione, contrariamente al caso dei reali, i punti del piano sono in numero finito, così le rette e le curve, come anche finito è il numero dei punti di una retta o di una curva.

Si ha così una geometria finita che trova numerosissime applicazioni anche perchè permette, per la risoluzione di molti problemi, l'uso del computer.

Ad esempio ci si può riferire al campo dei resti modulo p , che si indica con Z_p , e cioè al campo i cui elementi sono soltanto $0, 1, 2, \dots, p-1$ ed in cui $p=0$, $p+1=1$, $p+2=2, \dots$. Si dimostra che in questo caso, se p è primo, o la potenza di un numero primo $q=p^h$, esiste il piano affine in cui le coordinate dei punti e i coefficienti delle curve siano elementi di Z_p . Tale piano si indica con $AG(2, q)$.

* Lavoro fatto presso la sezione di Matematica del Dipartimento di Meccanica della Facoltà di Ingegneria di Roma III.

Il presente articolo è stato presentato, in un seminario, a docenti di Scuola Media Superiore.

Si dimostra che in $AG(2,q)$:

- il numero dei punti q^2 ;
- il numero delle rette è q^2+q ;
- il numero dei punti di una retta è q ;
- per ogni punto, passano $q+1$ rette.

Inoltre per le coniche irriducibili di $AG(2,q)$ si prova che:

- un'ellisse ha $q+1$ punti in $AG(2,q)$;
- una parabola ha q punti in $AG(2,q)$ (in quanto uno è all'infinito);
- un'iperbole ha $q-1$ punti in $AG(2,q)$ (in quanto due sono all'infinito);

Le cose vanno in modo diverso a seconda che q sia pari o dispari, come faremo vedere trattando qualche caso particolare.

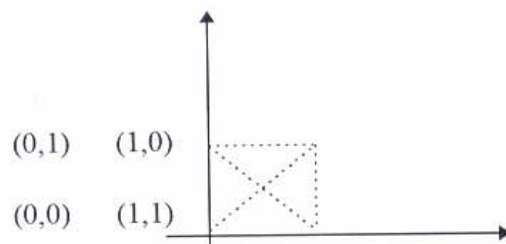
Considereremo i casi più semplici e precisamente i casi $q=2$ e $q=3$.

IL PIANO AFFINE $AG(2,2)$

Il piano affine costruito su Z_2 (i cui elementi sono solo 0 ed 1 ed in cui $1=-1$ e $2=0$), cioè $AG(2,2)$ è tale che:

- il numero dei punti è $q^2=4$;
- il numero delle rette è $q^2+q=6$;
- il numero dei punti di una retta è $q=2$ punti;
- per ogni punto passano $q+1=3$ rette.

I punti sono $(0,0)$, $(0,1)$, $(1,0)$ e $(1,1)$ e le rette $x=0$, $y=0$, $x+1=0$, $y+1=0$, $x+y=0$ e $x+y+1=0$.



Si noti che le rette sono costituite da soli due punti. Precisamente

- $x=0$ è costituita da $(0,0)$ e $(0,1)$;
- $y=0$ è costituita da $(0,0)$ e $(1,0)$;
- $x+1=0$ è costituita da $(1,0)$ e $(1,1)$;
- $y+1=0$ è costituita da $(0,1)$ e $(1,1)$;
- $x+y=0$ è costituita da $(0,0)$ e $(1,1)$;
- $x+y+1=0$ è costituita da $(1,0)$ e $(0,1)$.

Si verifica subito che per ogni punto passano tre rette (per (0,0), ad esempio, passano le tre rette $x=0$, $y=0$, $x+y=0$ e così via).

Vediamo come si comportano le coniche in $AG(2,2)$. Le cose si discostano dal caso q dispari e valgono anche in $AG(2,q)$ con $q=2^h$.

L'equazione di una conica in $AG(2,2)$ si scrive:

$$a_{11}x^2 + a_{22}y^2 + a_{12}xy + a_{13}x + a_{23}y + a_{33} = 0 \text{ dove } a_{ik} \in Z_2.$$

Se $a_{12}=a_{13}=a_{23}=0$ si ha:

$$a_{11}x^2 + a_{22}y^2 + a_{33} = 0$$

e cioè, poichè in Z_2 ogni elemento è un quadrato, ne segue:

$$(\sqrt{a_{11}}x + \sqrt{a_{22}}y + \sqrt{a_{33}})^2 = 0$$

cioè ogni conica è doppiamente degenere cioè è costituita da due rette coincidenti.

Si dimostra facilmente che vale anche il viceversa cioè *se la conica è doppiamente degenere si ha che $a_{12}=a_{13}=a_{23}=0$* .

Una conica è *semplicemente degenere* se è costituita da due rette distinte cioè se il polinomio di secondo grado che la rappresenta è il prodotto di due polinomi di primo grado distinti a coefficienti in Z_2 (o in un ampliamento quadratico di esso).

In $AG(2,2)$ i 6 coefficienti dell'equazione di una conica possono prendere solo i valori 0 ed 1 e quindi il numero di tali coniche sarà dato dalle disposizioni con ripetizione di due elementi a sei a sei. Da tale numero bisogna togliere la sestupla costituita da tutti zeri e quella in cui sono zero i primi cinque coefficienti e le sei che risultano di primo grado, quindi il numero sarà $2^6 - 8 = 56$.

Si possono avere mediante il seguente programma le equazioni di tutte le coniche di $AG(2,2)$.

```
{QUESTO PROGRAMMA SCRIVE LE EQUAZIONI DELLE}
{CONICHE}
```

```
PROGRAM coniche;
```

```
USES crt;
```

```
VAR i:INTEGER;
```

```
FUNCTION eq(index:BYTE):STRING;
```

```
VAR
```

```
  eqtmp : STRING;
```

```
  i : INTEGER;
```

```
BEGIN { eq }
```

```
  eqtmp:="";
```

```
  FOR i:=1 TO 6 DO
```

```
  BEGIN
```

```
    IF (index AND 1)=1 THEN
```

```
      CASE i OF
```

```
        1: eqtmp:='1'+eqtmp;
```

```

2: BEGIN
  IF (LENGTH(eqttmp)>0) THEN eqtmp:='++eqtmp;
    eqtmp:='y'+eqtmp
  END;
3: BEGIN
  IF (LENGTH(eqttmp)>0) THEN eqtmp:='++eqtmp;
    eqtmp:='x'+eqtmp
  END;
4: BEGIN
  IF (LENGTH(eqttmp)>0) THEN eqtmp:='++eqtmp;
    eqtmp:='xy'+eqtmp
  END;
5: BEGIN
  IF (LENGTH(eqttmp)>0) THEN eqtmp:='++eqtmp;
    eqtmp:='yy'+eqtmp
  END;
6: BEGIN
  IF (LENGTH(eqttmp)>0) THEN eqtmp:='++eqtmp;
    eqtmp:='xx'+eqtmp
  END
  END;
  index:=index SHR 1
  END;
  IF (LENGTH(eqttmp)>0) THEN eqtmp:=eqtmp+'=0';
  eq:=eqtmp;
END; { eq }
BEGIN { main }
  FOR i:=8 TO 63 DO
    WRITE(eq(i):20);
  READLN;
END. { main }

```

Come equazioni si avranno in output:

- | | |
|------------------|--------------------|
| 1) $xy=0$ | 2) $xy+1=0$ |
| 3) $xy+y=0$ | 4) $xy+y+1=0$ |
| 5) $xy+x=0$ | 6) $xy+x+1=0$ |
| 7) $xy+x+y=0$ | 8) $xy+x+y+1=0$ |
| 9) $y^2=0$ | 10) $y^2+1=0$ |
| 11) $y^2+y=0$ | 12) $y^2+y+1=0$ |
| 13) $y^2+x=0$ | 14) $y^2+x+1=0$ |
| 15) $y^2+x+y=0$ | 16) $y^2+x+y+1=0$ |
| 17) $y^2+xy=0$ | 18) $y^2+xy+1=0$ |
| 19) $y^2+xy+y=0$ | 20) $y^2+xy+y+1=0$ |
| 21) $y^2+xy+x=0$ | 22) $y^2+xy+x+1=0$ |

- | | |
|------------------------|--------------------------|
| 23) $y^2+xy+x+y=0$ | 24) $y^2+xy+x+y+1=0$ |
| 25) $x^2=0$ | 26) $x^2+1=0$ |
| 27) $x^2+y=0$ | 28) $x^2+y+1=0$ |
| 29) $x^2+x=0$ | 30) $x^2+x+1=0$ |
| 31) $x^2+x+y=0$ | 32) $x^2+x+y+1=0$ |
| 33) $x^2+xy=0$ | 34) $x^2+xy+1=0$ |
| 35) $x^2+xy+y=0$ | 36) $x^2+xy+y+1=0$ |
| 37) $x^2+xy+x=0$ | 38) $x^2+xy+x+1=0$ |
| 39) $x^2+xy+x+y=0$ | 40) $x^2+xy+x+y+1=0$ |
| 41) $x^2+y^2=0$ | 42) $x^2+y^2+1=0$ |
| 43) $x^2+y^2+y=0$ | 44) $x^2+y^2+y+1=0$ |
| 45) $x^2+y^2+x=0$ | 46) $x^2+y^2+x+1=0$ |
| 47) $x^2+y^2+x+y=0$ | 48) $x^2+y^2+x+y+1=0$ |
| 49) $x^2+y^2+xy=0$ | 50) $x^2+y^2+xy+1=0$ |
| 51) $x^2+y^2+xy+y=0$ | 52) $x^2+y^2+xy+y+1=0$ |
| 53) $x^2+y^2+xy+x=0$ | 54) $x^2+y^2+xy+x+1=0$ |
| 55) $x^2+y^2+xy+x+y=0$ | 56) $x^2+y^2+xy+x+y+1=0$ |

Per vedere quali di tali coniche sono ellissi, parabole, iperboli non si può seguire il criterio che si applica nel caso del piano affine sui reali in cui l'equazione di una conica si scrive

$$a_{11}x^2+a_{22}y^2+2a_{12}xy+2a_{13}x+2a_{23}y+a_{33}=0 \text{ dove } a_{ik} \in \mathbb{R}$$

ed in cui si vede se $\Delta = a_{12}^2 - a_{11}a_{22} \geq < 0$ (dove Δ è il discriminante dell'equazione che si ottiene ponendo uguali a zero i termini di 2° grado della suddetta equazione), poichè in $\mathbb{Z}_2$ non si può applicare la formula risolutiva dell'equazione di 2° grado, essendo $2=0$.

Basta però ricordare (cfr. Introduzione) che: in $AG(2,2)$ un'ellisse ha tre punti, una parabola ne ha due e un'iperbole ne ha uno.

Ad esempio la conica $x^2+y^2+xy+1=0$ contiene i punti (1,0), (0,1), (1,1) e quindi è un'ellisse; la conica $x+y^2+1=0$ contiene i punti (0,1) e (1,0) e quindi è una parabola; la conica $xy+x+y=0$ contiene il punto (0,0) e quindi è un'iperbole.

Si noti ancora come esempio che la conica $x^2+y^2+1=0$ risulta doppiamente degenerare $(x+y+1)^2=0$ e la conica $x^2+xy=0$ risulta semplicemente degenerare in quanto si spezza nelle rette $x=0$ e $x+y=0$.

Si può scrivere il seguente programma che riconosce se una conica è degenerare (cioè si spezza nel prodotto di due delle sei rette di $AG(2,2)$, distinte o coincidenti) e nel caso in cui sia non degenerare dice se è un'ellisse, una parabola, un'iperbole.

```
{QUESTO PROGRAMMA RICONOSCE SE UNA CONICA E'}  
{DEGENERARE E, IN TAL CASO, TROVA LE EQUAZIONI DELLE}  
{RETTE IN CUI SI SPEZZA E, NEL CASO IN CUI NON E'}  
{DEGENERARE, DICE SE E' UN'ELLISSE, UNA PARABOLA OPPURE}  
{UN'IPERBOLE.}
```

```
PROGRAM riconoscitore;  
USES crt;  
FUNCTION eq(index:BYTE):STRING;  
VAR  
  eqtmp : STRING;  
  i      : INTEGER;  
BEGIN { eq }  
  eqtmp:="";  
  FOR i:=1 TO 6 DO  
    BEGIN  
      IF (index AND 1)=1 THEN  
        CASE i OF  
          1: eqtmp:='1'+eqtmp;  
          2: BEGIN  
              IF (LENGTH(eqtmp)>0) THEN eqtmp:='+'+eqtmp;  
              eqtmp:='y'+eqtmp  
            END;  
          3: BEGIN  
              IF (LENGTH(eqtmp)>0) THEN eqtmp:='+'+eqtmp;  
              eqtmp:='x'+eqtmp  
            END;  
          4: BEGIN  
              IF (LENGTH(eqtmp)>0) THEN eqtmp:='+'+eqtmp;  
              eqtmp:='xy'+eqtmp  
            END;  
          5: BEGIN  
              IF (LENGTH(eqtmp)>0) THEN eqtmp:='+'+eqtmp;  
              eqtmp:='yy'+eqtmp  
            END;  
          6: BEGIN  
              IF (LENGTH(eqtmp)>0) THEN eqtmp:='+'+eqtmp;  
              eqtmp:='xx'+eqtmp  
            END  
          END;  
      index:=index SHR 1  
    END;  
  IF (LENGTH(eqtmp)>0) THEN eqtmp:=eqtmp+'0';  
  eq:=eqtmp;
```

```

END; { eq }
FUNCTION eqprod(eq1,eq2:BYTE):BYTE;
VAR
  m,n,o,p,q,r : BYTE;
BEGIN { eqprod }
  m:=((eq1 SHR 2) AND (eq2 SHR 2)) AND 1;
  n:=((eq1 SHR 1) AND (eq2 SHR 1)) AND 1;
  r:=(eq1 AND eq2) AND 1;
  o:=((eq1 SHR 2) AND (eq2 SHR 1) AND 1)+((eq1 SHR 1) AND (eq2 SHR
2) AND 1);
  p:=((eq1 SHR 2) AND eq2 AND 1)+(eq1 AND (eq2 SHR 2) AND 1);
  q:=((eq1 SHR 1) AND eq2 AND 1)+(eq1 AND (eq2 SHR 1) AND 1);
  o:=o MOD 2;
  p:=p MOD 2;
  q:=q MOD 2;
  eqprod:=(r AND 1) OR ((q SHL 1) AND 2) OR
              ((p SHL 2) AND 4) OR
              ((o SHL 3) AND 8) OR
              ((n SHL 4) AND 16) OR
              ((m SHL 5) AND 32)
END; { eqprod }
FUNCTION evaluate(index,x,y:BYTE):BOOLEAN;
VAR
  v : BYTE;
BEGIN { evaluate }
  IF (x<>0) THEN x:=1;
  IF (y<>0) THEN y:=1;
  v:= (((index SHR 5) AND 1)*x) +
      (((index SHR 4) AND 1)*y) +
      (((index SHR 3) AND 1)*x*y) +
      (((index SHR 2) AND 1)*x) +
      (((index SHR 1) AND 1)*y) + (index AND 1)) MOD 2;
  IF (v=0) THEN evaluate:=TRUE
  ELSE evaluate:=FALSE
END; { evaluate }
FUNCTION recognize(index:BYTE):STRING;
VAR
  i,j,num : BYTE;
BEGIN { recognize }
  recognize:="";
  num:=0;
  FOR i:=0 TO 1 DO
    FOR j:=0 TO 1 DO
      IF evaluate(index,i,j) THEN num:=num+1;

```

```

CASE num OF
  1: recognize:='Iperbole';
  2: recognize:='Parabola';
  3: recognize:='Ellisse'
END
END; { recognize }
PROCEDURE exclude;
CONST
  rette : ARRAY [1..6] OF BYTE = (4,2,5,3,6,7);
VAR
  i,j,k : BYTE;
  flag : BOOLEAN;
  f : TEXT;
BEGIN { exclude }
  ASSIGN(f,'r.asc');
  REWRITE(f);
  FOR i:=8 TO 63 DO
    BEGIN
      WRITE(f,eq(i):20);
      flag:=TRUE;
      FOR j:=1 TO 6 DO
        FOR k:=j TO 6 DO
          IF (eqprod(rette[j],rette[k])=i) THEN
            BEGIN
              flag:=FALSE;
              WRITELN(f,'      [' ,eq(rette[j]),'] * [' ,eq(rette[k]),']')
            END;
          IF flag THEN WRITELN(f,'      ',recognize(i))
        END;
      CLOSE(f)
    END; { exclude }
  BEGIN { main }
  exclude
END. { main }

```

In output si hanno i seguenti risultati:

- | | |
|-----------------|-------------------|
| 1) $xy=0$ | $[x=0]*[y=0]$ |
| 2) $xy+1=0$ | Iperbole |
| 3) $xy+y=0$ | $[y=0]*[x+1=0]$ |
| 4) $xy+y+1=0$ | Iperbole |
| 5) $xy+x=0$ | $[x=0]*[y+1=0]$ |
| 6) $xy+x+1=0$ | Iperbole |
| 7) $xy+x+y=0$ | Iperbole |
| 8) $xy+x+y+1=0$ | $[x+1=0]*[y+1=0]$ |

9) $y^2=0$	$[y=0]*[y=0]$
10) $y^2+1=0$	$[y+1=0]*[y+1=0]$
11) $y^2+y=0$	$[y=0]*[y+1=0]$
12) $y^2+y+1=0$	
13) $y^2+x=0$	Parabola
14) $y^2+x+1=0$	Parabola
15) $y^2+x+y=0$	Parabola
16) $y^2+x+y+1=0$	Parabola
17) $y^2+xy=0$	$[y=0]*[x+y=0]$
18) $y^2+xy+1=0$	Iperbole
19) $y^2+xy+y=0$	$[y=0]*[x+y+1=0]$
20) $y^2+xy+y+1=0$	Iperbole
21) $y^2+xy+x=0$	Iperbole
22) $y^2+xy+x+1=0$	$[y+1=0]*[x+y+1=0]$
23) $y^2+xy+x+y=0$	$[y+1=0]*[x+y=0]$
24) $y^2+xy+x+y+1=0$	Iperbole
25) $x^2=0$	$[x=0]*[x=0]$
26) $x^2+1=0$	$[x+1=0]*[x+1=0]$
27) $x^2+y=0$	Parabola
28) $x^2+y+1=0$	Parabola
29) $x^2+x=0$	$[x=0]*[x+1=0]$
30) $x^2+x+1=0$	
31) $x^2+x+y=0$	Parabola
32) $x^2+x+y+1=0$	Parabola
33) $x^2+xy=0$	$[x=0]*[x+y=0]$
34) $x^2+xy+1=0$	Iperbole
35) $x^2+xy+y=0$	Iperbole
36) $x^2+xy+y+1=0$	$[x+1=0]*[x+y+1=0]$
37) $x^2+xy+x=0$	$[x=0]*[x+y+1=0]$
38) $x^2+xy+x+1=0$	Iperbole
39) $x^2+xy+x+y=0$	$[x+1=0]*[x+y=0]$
40) $x^2+xy+x+y+1=0$	Iperbole
41) $x^2+y^2=0$	$[x+y=0]*[x+y=0]$
42) $x^2+y^2+1=0$	$[x+y+1=0]*[x+y+1=0]$
43) $x^2+y^2+y=0$	Parabola
44) $x^2+y^2+y+1=0$	Parabola
45) $x^2+y^2+x=0$	Parabola
46) $x^2+y^2+x+1=0$	Parabola
47) $x^2+y^2+x+y=0$	$[x+y=0]*[x+y+1=0]$
48) $x^2+y^2+x+y+1=0$	
49) $x^2+y^2+xy=0$	Iperbole
50) $x^2+y^2+xy+1=0$	Ellisse
51) $x^2+y^2+xy+y=0$	Ellisse
52) $x^2+y^2+xy+y+1=0$	Iperbole

- 53) $x^2+y^2+xy+x=0$ Ellisse
 54) $x^2+y^2+xy+x+1=0$ Iperbole
 55) $x^2+y^2+xy+x+y=0$ Ellisse
 56) $x^2+y^2+xy+x+y+1=0$ Iperbole

Si noti che il programma non ha specificato tre coniche, perchè esse sono degeneri, ma si spezzano in rette che non appartengono al campo Z_2 , in quanto i polinomi che le rappresentano sono irriducibili in Z_2 , ma riducibili in un ampliamento algebrico (come succede per i polinomi che sono irriducibili nel campo dei reali, ma riducibili nel campo dei complessi).

IL PIANO AFFINE $AG(2,3)$

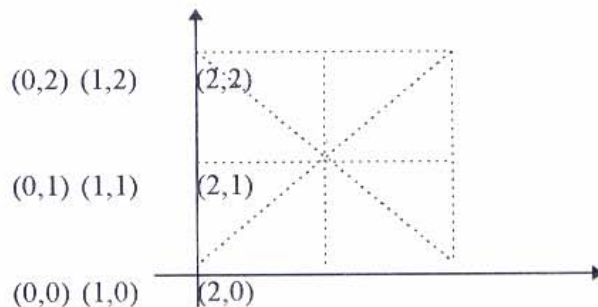
Consideriamo ora $AG(2,3)$ cioè il piano affine costruito su Z_3 e cioè sul campo dei resti modulo 3, i cui elementi sono quindi solo 0, 1, 2 (ed in cui $3=0$, $1=-2$).

In $AG(3,q)$:

- i punti sono in numero di $q^2=3^2=9$;
- le rette sono in un numero di $q^2+q=3^2+3=12$;
- il numero di punti di una retta sono $q=3$;
- per un punto passano quattro rette.

I punti sono:

$(0,0)$, $(1,0)$, $(2,0)$, $(0,1)$, $(0,2)$, $(1,1)$, $(2,2)$, $(1,2)$, $(2,1)$.



Le 12 rette (ognuna delle quali contiene 3 punti) si ottengono dall'equazione $ax+by+c=0$ prendendo i coefficienti appartenenti a Z_3 e togliendo una retta se ha un coefficiente uguale a 2 e si può ottenere da un'altra quando si moltiplicano i coefficienti per 2 e si prendono modulo 3. Ad esempio $x+2y=0$ si può scrivere $2x+4y=0$ cioè $2x+y=0$ ed essendo tali rette coincidenti, bisogna eliminarne una.

Possiamo allora dire che le 12 rette hanno per equazioni le seguenti:

- $x=0$ che è costituita da $(0,0)$, $(0,1)$ e $(0,2)$;
 $y=0$ che è costituita da $(0,0)$, $(1,0)$ e $(2,0)$;

$x+y=0$	che è costituita da (0,0), (1,2) e (2,1);
$2x+y=0$	che è costituita da (0,0), (1,1) e (2,2);
$x+y+1=0$	che è costituita da (1,1), (2,0) e (0,2);
$2x+2y+1=0$	che è costituita da (1,0), (0,1) e (2,2);
$2x+1=0$	che è costituita da (1,0), (1,1) e (1,2);
$x+1=0$	che è costituita da (2,0), (2,2) e (2,1);
$y+1=0$	che è costituita da (0,2), (2,2) e (1,2);
$2x+y+1=0$	che è costituita da (0,2), (1,0) e (2,1);
$2x+y+2=0$	che è costituita da (2,0), (0,1) e (1,2);
$2y+1=0$	che è costituita da (0,1), (1,1) e (2,1).

Per ogni punto passano 4 rette, come subito si verifica (ad esempio per 0 passano le rette $x=0$, $y=0$, $x+y=0$, $2x+y=0$).

Si possono scrivere le equazioni suddette tramite il seguente programma.

{QUESTO PROGRAMMA SCRIVE LE EQUAZIONI DELLE RETTE}

PROGRAM rette;

USES crt;

VAR

eq : ARRAY [1..27] OF BOOLEAN;

PROCEDURE eliminal;

VAR

a,b,c,i : BYTE;

BEGIN { eliminal }

i:=1;

FOR a:=0 TO 2 DO

FOR b:=0 TO 2 DO

FOR c:=0 TO 2 DO

BEGIN

eq[i]:=TRUE;

IF ((a=0) AND (b=0)) THEN eq[i]:=FALSE;

IF ((a=2) AND (b=2) AND (c=2)) THEN eq[i]:=FALSE;

IF (a=0) THEN

IF ((b=2) AND (c=2)) THEN eq[i]:=FALSE;

IF (b=0) THEN

IF ((a=2) AND (c=2)) THEN eq[i]:=FALSE;

IF (c=0) THEN

IF ((a=2) AND (b=2)) THEN eq[i]:=FALSE;

IF (a=2) THEN

IF ((b=0) AND (c=0)) THEN eq[i]:=FALSE;

IF (b=2) THEN

IF ((a=0) AND (c=0)) THEN eq[i]:=FALSE;

IF (c=2) THEN

IF ((a=0) AND (b=0)) THEN eq[i]:=FALSE;

```

        i:=i+1
    END
END; { elimina1 }
PROCEDURE elimina2;
VAR
    a,b,c,aa,bb,cc,i : BYTE;
BEGIN { elimina2 }
    i:=1;
    FOR a:=0 TO 2 DO
        FOR b:=0 TO 2 DO
            FOR c:=0 TO 2 DO
                BEGIN
                    IF eq[i] THEN
                        IF ((a=2) OR (b=2) OR (c=2)) THEN
                            BEGIN
                                aa:=(a*2) MOD 3;
                                bb:=(b*2) MOD 3;
                                cc:=(c*2) MOD 3;
                                IF eq[(aa*9)+(bb*3)+cc+1] THEN eq[i]:=FALSE
                            END;
                        i:=i+1
                    END
                END
            END; { elimina2 }
PROCEDURE stampa;
VAR
    a,b,c,i : BYTE;
    f      : TEXT;
BEGIN { stampa }
    ASSIGN(f,'r1.asc');
    REWRITE(f);
    i:=1;
    FOR a:=0 TO 2 DO
        FOR b:=0 TO 2 DO
            FOR c:=0 TO 2 DO
                BEGIN
                    IF eq[i] THEN
                        BEGIN
                            IF (a<>0) THEN
                                IF (a=1) THEN WRITE(f,'x')
                                    ELSE WRITE(f,'2x');
                                IF (b<>0) THEN
                                    BEGIN
                                        IF (a<>0) THEN WRITE(f,'+');
                                        IF (b=1) THEN WRITE(f,'y')

```

```

ELSE WRITE(f,'2y');
END;
IF (c<>0) THEN WRITE(f,'+',c);
WRITELN(f,'=0')
END;
i:=i+1
END;
CLOSE(f)
END; { stampa }
BEGIN { main }
  elimina1;
  elimina2;
  stampa
END. { main }

```

Per quanto riguarda le coniche esse si possono scrivere come nel caso del piano affine su $\mathbb{R}$:

$$a_{11}x^2 + a_{22}y^2 + 2a_{12}xy + 2a_{13}x + 2a_{23}y + a_{33} = 0$$

in cui gli $a_{ik} \in \mathbb{Z}_3$ cioè possono essere 0, 1, 2.

All'uopo si potrebbero scrivere le equazioni di tutte le coniche con un opportuno programma.

Una conica è degenere se

$$\Delta = \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{12} & a_{22} & a_{23} \\ a_{13} & a_{23} & a_{33} \end{vmatrix} = 0$$

ed è rispettivamente una iperbole, parabola, ellisse a seconda che, calcolando

$$-\begin{vmatrix} a_{11} & a_{12} \\ a_{12} & a_{22} \end{vmatrix} \text{ si ha:}$$

$$a_{12}^2 - a_{11}a_{22} = \square \text{ quadrato nel campo cioè } 1$$

$$a_{12}^2 - a_{11}a_{22} = 0$$

$$a_{12}^2 - a_{11}a_{22} = \square \text{ non quadrato nel campo cioè } 2$$

In $AG(2,3)$ un'ellisse ha $q+1=4$ punti, quelli di una parabola sono $q=3$, quelli di un'iperbole $q-1=2$.

Esempi:

La conica $2x^2 + y^2 + 2xy + 1 = 0$, poichè $a_{12}^2 - a_{11}a_{22} = 1 - 2 = 1 - 2 = -1 = 2$, è un'ellisse e contiene i punti (0,1), (1,1), (2,2), (2,0).

La conica $x^2 + y^2 + 2xy + x = 0$, poichè $a_{12}^2 - a_{11}a_{22} = 0$, è una parabola e contiene i punti (0,0), (2,0), (2,2).

La conica $2x^2 + y^2 + x = 0$, poichè $a_{12}^2 - a_{11}a_{22} = -2 = 1$, è un'iperbole contiene i punti (0,0), (1,0).

La conica $x^2+y^2+2xy+x+y=0$ è semplicemente degenere in quanto (si noti che $1=4$, cioè $\frac{1}{2}=2$)

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

e si scrive $(x+y)^2+(x+y)=0$ $(x+y)(x+y+1)=0$.

La conica $x^2+y^2+2xy+2x+2y+1=0$ è doppiamente degenere e si scrive $(x+y+2)^2=0$.

Esistono programmi per il calcolo del determinante di una matrice e quindi, mediante essi, si può riconoscere se una conica è degenere e, in caso contrario, se è un'ellisse, una parabola o un'iperbole.

BIBLIOGRAFIA

1. B. SEGRE, *Lectures on modern geometry*, Edizioni Cremonese, Roma 1961.

IL VANTAGGIO DI OPERARE IN UN AMBIENTE FINITO

Mauro Zannetti *

1. Introduzione

La fattorizzazione di un polinomio e la determinazione dei suoi fattori lineari, è un argomento che negli ultimi anni ha riscontrato un grande interesse rivolto alla ricerca di algoritmi di risoluzione di bassa complessità computazionale. Nel presente articolo, testo di una lezione tenuta dall'autore presso la sezione provinciale Mathesis di Teramo nel dicembre del 1995 viene trattato il calcolo degli zeri razionali di un polinomio, tenendo conto che ciò equivale alla determinazione dei fattori lineari. La storia del problema della fattorizzazione di un polinomio, è lunga ed interessante. I primi algoritmi, per trovare i fattori lineari di un polinomio con coefficienti interi, furono presentati da Isaac Newton nel 1707, e dall'astronomo Friederich T.V. Schubert nel 1793. Un importante criterio per determinare l'irriducibilità di un polinomio in $Z[x]$ fu dato da F.C. Eisenstein nel 1850.

L. Kronecker riscoprì il metodo di Schubert nel 1882 e trovò un algoritmo per fattorizzare polinomi con due o più variabili o con coefficienti in estensioni algebriche. Purtroppo quando, circa cento anni più tardi, questi algoritmi furono implementati sul calcolatore, essi risultarono essere altamente inefficienti. Nel 1967 E. Berlekamp, scoprì un ingegnoso algoritmo che permise di fattorizzare sul campo $Z_p[x]$, p numero primo, il cui tempo di calcolo risultò notevolmente minore dei precedenti. La tecnica, interessante di per sé, mette in risalto il vantaggio di trasportare un problema in un ambiente finito, risolverlo e poi tornare nell'ambiente di provenienza. Introduciamo ora il problema della fattorizzazione lineare.

Dato un polinomio $A(x) = \sum_{i=0}^n a_i x^i$ di grado n con coefficienti di dimensione arbitraria, $a_i \in Z$ per $i = 0, 1, \dots, n$, determiniamo, i numeri

* Liceo Scientifico "V. Pollione" Avezzano (AQ).

razionali s/t , $t \neq 0$ tali che $A(s/t) = 0$. Il problema è equivalente a trovare tutti i fattori lineari di $A(x)$; infatti, s/t è uno zero razionale, se e solo se $(tx-s)$ divide $A(x)$. Supponiamo che $A(x)$ sia un polinomio monico. Fattorizzare $A(x)$ significa esprimere $A(x)$ nella forma:

$$A(x) = p_1(x)^{e_1} \times p_2(x)^{e_2} \times \dots \times p_r(x)^{e_r},$$

dove $p_1(x), p_2(x), \dots, p_r(x)$ sono polinomi monici distinti irriducibili. Una tecnica standard per individuare i fattori di molteplicità maggiore di uno è calcolare il Massimo Comun Divisore tra $A(x)$ e la sua derivata $A'(x)$; sia esso uguale a $d(x)$.

Se $d(x)$ è uguale a 1, sappiamo che $A(x)$ è privo di quadrati, infatti: se

$$A(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_0 = (V(x))^2 \times W(x) \text{ allora la sua derivata è}$$

$$A'(x) = n a_n x^{n-1} + (n-1) a_{n-1} x^{n-2} + \dots + a_1 =$$

$$V(x) [2V'(x)W(x) + V(x)W'(x)]$$

e $A(x)$ è il prodotto di polinomi primi $p_1(x), p_2(x), \dots, p_r(x)$. Se $d(x)$ non è uguale a 1, allora $d(x)$ è un fattore proprio di $A(x)$; e così il processo continua fattorizzando $d(x)$ e $A(x)/d(x) = A_1(x)$ separatamente. Si noti che $A_1(x)$ è ora sicuramente privo di quadrati, e su $d(x)$ dovrà essere eseguito lo stesso processo. L'algoritmo di Newton sfrutta il seguente teorema.

Teorema. Sia $A(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_0$ un polinomio in $Z[x]$ e $x = s/t$ una sua radice, con s, t interi relativamente primi, allora s/a_0 e t/a_n .

Dim.: Per ipotesi si ha $A(s/t) = 0$ allora $a_n (s^n/t^n) + a_{n-1} (s^{n-1}/t^{n-1}) + \dots + a_1 (s/t) + a_0 = 0$, moltiplicando tutto per t^n si ha: $a_n s^n = -t [a_{n-1} (s^{n-1}) + \dots + a_1 (s t^{n-2}) + a_0 t^{n-1}]$ per cui t deve dividere $a_n s^n$ e poichè s, t sono relativamente primi, t deve dividere a_n . Ugualmente s deve dividere $a_0 t^n$ e per la stessa ragione s deve dividere a_0 .

L'algoritmo è basato sulla ricerca esaustiva tra tutti i fattori di a_0 , a_n che potrebbero essere eventuali soluzioni di $A(x) = 0$.

Il metodo, ovviamente, è utile solamente per polinomi con a_0 e a_n , interi con poche cifre, per cui ridefiniamo il problema in un dominio immagine con elementi di dimensione fissata. Consideriamo Z_p , gli interi modulo p , dove p è un numero primo.

2. Fattorizzazione lineare su $Z_p[x]$

Trasportiamo $A(x)$ in $Z_p[x]$ considerando $A^*(x) \equiv A(x) \pmod{p}$ con coefficienti in Z_p . Appliciamo ora la ricerca esaustiva ai p elementi di Z_p , se $A^*(p_i) = 0$ ($p_i \in \{0, 1, \dots, p-1\}$) allora p_i è uno zero di $A^*(x)$. In generale $A^*(x)$ avrà più zeri di $A(x)$ infatti, ad esempio, il polinomio $X^4 + 1$ non ha zeri in Z , ma ha uno zero in Z_2 il problema è risalire alle soluzioni in Z partendo dalle soluzioni in Z_p . Le tecniche più utilizzate sono due: quella classica dei Resti Cinesi e quella più moderna del lifting.

Lifting lineare

Teorema. Sia $A(x)$ un polinomio a coefficienti in Z e sia a uno zero semplice di $A(x) \pmod{p}$ con p primo tale che a non è uno zero di $A'(x) \pmod{p}$. Allora $\forall r \geq 1$ l'equazione $A(x) \equiv 0 \pmod{p^r}$ ha una soluzione $a_r \equiv a \pmod{p}$.

Dim: procediamo per induzione: per $r=1$ il teorema è vero per ipotesi. Supponiamo vera l'ipotesi per $r \leq n$. Sia $b_n = A(a_n) / p^n$ la divisione è esatta per l'ipotesi di induzione e sia $a_{n+1} = a_n + p^n y$ con y indeterminata.

$$A(a_{n+1}) \equiv \left(A(a_n) + p^n y A'(a_n) \right) \pmod{p^{n+1}} \equiv$$

Allora

$$p^n \left(b_n + A'(a) y \right) \pmod{p^{n+1}}.$$

Dal momento che per ipotesi si ha $A'(a) \equiv 0 \pmod{p}$, possiamo scegliere $y = \left(-b_n / A'(a) \right) \pmod{p}$ ottenendo $A(a_{n+1}) \equiv 0 \pmod{p^{n+1}}$. Per cui il passo di lifting sarà:

$$a_{n+1} = a_n - p^n \left(\left(A(a_n) / p^n \right) / A'(a) \right) \pmod{p}.$$

Lifting Quadratico

E' una modifica del lifting lineare. Il teorema che segue, di cui omettiamo la dimostrazione, illustra tale metodo che è più veloce di quello lineare, in quanto sono necessari meno passi per raggiungere uno stesso limite.

Teorema. Sia $A(x)$ un polinomio a coefficienti in Z e a uno zero di $A(x)(\text{mod } p)$ con p primo tale che a non sia uno zero di $A'(x)(\text{mod } p)$. Allora $\forall r \geq 1$ l'equazione $A(x) \equiv 0 \pmod{p^{2^r}}$ ha una soluzione $a_{2^r} \equiv a(\text{mod } p)$.

3. Polinomi non monici

Rimuoviamo ora la condizione di monicità sul polinomio. Consideriamo il polinomio $A(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_0$ con coefficienti in Z e scegliamo un primo p tale che p non divide a_n e consideriamo $A_1(x) \equiv A(x)/a_n \pmod{p}$. Sia a_1 uno zero di $A_1(x)$; utilizziamo le tecniche di lifting ricaviamo a_n^* il corrispondente zero di $A_n(x) \pmod{p^n}$. Il seguente teorema illustra il legame tra a_n^* e la corrispondente soluzione del tipo s/t di $A(x)$.

Teorema. Sia $A(x)$ un polinomio a coefficienti in Z , p un primo tale che p non divide a_n e n un intero positivo. Se $A(x)$ ha uno zero razionale $b = s/t$ allora il polinomio $A_N(x) \equiv A(x)/a_n \pmod{p^N}$ ha uno zero b_N in Z_{p^N} e $b_N \equiv s \bar{a}_n a_n^{-1} \pmod{p^N}$, dove $\bar{a}_n$ è l'intero a_n/t .

Dim: Poichè $A(s/t) = 0$ segue che $a_n = \bar{a}_n \times t$. Siccome per ipotesi p non divide a_n , segue che p non divide t e quindi sia a_n che t hanno inversi in Z_{p^N} per ogni N maggiore di zero; pertanto

$$b_N \equiv s \times t^{-1} \equiv s \times t^{*1} \times a_n \times a_n^{-1} = s \times \bar{a}_n \times a_n^{-1} \pmod{p^N}.$$

poichè l'inverso di questo teorema è falso, dobbiamo sempre verificare che la soluzione s/t , ottenuta da b_N ponendo $s = b_n \times a_n \pmod{p^N}$ e $t = a_n \pmod{p^N}$ sia effettivamente uno zero di $A(x)$.

4. Algoritmo lineare di Newton-Hensel-Loos

Sia dato un polinomio privo di zeri multipli $A(x) \in \mathbb{Z}[x]$; il seguente algoritmo restituisce una lista L dei suoi zeri razionali;

- 1) cercare il più piccolo primo p tale che $A(x) \pmod{p}$ è privo di quadrati;
- 2) sia $b = 2|a_0||a_n|$ e scegliamo il più piccolo N tale che b sia minore di p^N ;
- 3) calcoliamo la lista L^* degli zeri di $a_n^{-1}A(x) \pmod{p}$ mediante l'algoritmo di Newton-Hensel;
- 4) eleviamo gli zeri in L^* di $a_n^{-1}A(x) \pmod{p}$ agli zeri di $a_n^{-1}A(x) \pmod{p^N}$ (lifting lineare);
- 5) per ogni a^* di L^* riduciamo $(a^* a_n \pmod{p^N}) / a_n \pmod{p}$ al termine s/t e se $A(s/t) = 0$, aggiungi s/t a L .

Come miglioramento dell'algoritmo precedente possiamo modificare il passo 4) sostituendo l'algoritmo di lifting lineare con quello quadratico. Gli algoritmi sono stati implementati su sistema Vax / 750, presso il Centro di calcolo dell'Università degli Studi de L'Aquila. L'implementazione è stata effettuata utilizzando il linguaggio Macsyma (Man and Computer Symbolic Manipulation System), linguaggio sviluppato presso il MIT, che è uno dei linguaggi di manipolazione algebrica più potenti tra quelli attualmente disponibili.

Qui di seguito riportiamo le tabelle con i tempi di esecuzione relative all'algoritmo lineare e a quello quadratico.

Algoritmo lineare

<i>grado dei polinomi</i>	<i>lunghezza delle soluzioni in cifre (tempi di esecuzione)</i>		
	15	20	40
2	1'	1' 40"	2' 7"
3	1' 2"	2' 1"	5' 6"
5	3' 3"	5' 1"	13' 2"
10	20' 2"	28' 8"	2h 0' 6"

Algoritmo quadratico

<i>grado dei polinomi</i>	<i>lunghezza delle soluzioni in cifre (tempi di esecuzione)</i>		
	15	20	40
2	31"	33"	46"
3	43"	47"	1' 3"
5	1' 10"	1' 12"	2' 3"
10	4' 4"	5' 12"	12' 5"
20	19' 1"	26' 13"	53' 3"

Dalle tabelle si può notare come l'algoritmo quadratico sia asintoticamente più efficiente di quello lineare. Ad esempio per polinomi di grado 20 con soluzioni di 40 cifre decimali il tempo di calcolo dell'algoritmo quadratico è circa la metà di quello lineare.

BIBLIOGRAFIA

- 1 R. BOGEN, *Macsyma Reference Manual*, The Matlab Group Laboratory For Computers Science M.I.T. Version 1983.
- 2 J. CALMET, R. LOOS, *Deterministic Versus Probabilistic Factorization of integral Polinomials* 1981.
- 3 D. S. HIRSCHBERG, C. k. WONG, *A polinomial Time Algorithm for the Knapsack Problem with two variables*, JACM 1976.
- 4 E. KALTOFEN, NEWARK, DELAWARE, *Factorization of polynomials*, Computing, Suppl. 1982.
- 5 D. E. KNUTH, *The Art of Computer Programming*, (vol. 2)- Addison Wesley 1969.
- 6 R. LOOS, *Computing Rationals Zeros of integral Polinomials by p -adic Expantion*-SIAM J. Computing 1983.
- 7 R. PAVELLE, M. ROTHSTEIN, J.FITCH, *Computer Algebra*, Scientific American 1981.
- 8 R. H. RAND, *Computer Algebra in Applied Mathematics*, An introduction to Macsyma-Research Notes in Mathematics n. 94, Pitman Advanced Publishing Program 1984.
- 9 D. Y. Y. YUNG, *Algebraic Algorithms using p -adic Construction*, ACM Symposium 1976.

BRUNO RIZZI E IL FUSIONISMO

Ciro D'Aniello *

Il presente non vuole essere un necrologio, ma semplicemente un commiato da un vero amico, sostenuto dalla speranza di tenerne vivo il ricordo negli amici e in quanti hanno avuto la fortuna di incontrarlo.

Ringrazio gli organizzatori del convegno, ed in particolare il Prof. F. Eugeni, per avermi invitato a partecipare e a dare il mio contributo consistente nel ricordare qualche aspetto della multiforme attività scientifica e didattica del compianto B. Rizzi.

In particolare, voglio ricordare l'attività didattica del Nostro che, sull'esempio del Prof. B. de Finetti, aderì a quella corrente della didattica della matematica chiamata "fusionismo" apportandovi contributi originali tramite scritti, articoli, conferenze.

Il fusionismo nasce con F. Klein e il suo Programma di Erlangen del 1872, con la condanna dello stesso Klein dello specialismo in matematica e la sua preferenza per l'intuizione; ne segue un approccio didattico di natura interdisciplinare, che fa superare la schematica e antiquata suddivisione della matematica in aritmetica, algebra, analisi, topologia, etc.

Come ha ricordato il Nostro nel suo ultimo articolo [1], dal titolo "de Finetti e il Periodico di Matematiche", il fusionismo nasce come "quella corrente della didattica che vuole principalmente la completa fusione degli argomenti geometrici e di quelli analitici (aritmetici, algebrici, etc.), tradizionalmente tenuti separati, spesso addirittura con bigottismo puristico quasi per timore che al contatto si contaminassero a vicenda". Ancora dallo stesso articolo: "l'insegnamento per problemi, nella concezione fusionista, si inserisce e nasce in modo naturale; non è qualcosa di artificioso, ma al contrario è qualcosa di intimamente connesso al modo di concepire la matematica e di volerla divulgare e insegnare".

* Preside Ist. Mag. St. "D. Falconi" Velletri (Roma).

Da queste premesse metodologiche ne conseguono necessariamente l'integrazione e la connessione tra le diverse parti della matematica, comprese le relative applicazioni; inoltre l'insegnamento della stessa deve essere condotto con il metodo socratico della maieutica, secondo il quale uno schiavo ignorante viene condotto a risolvere problemi.

Il punto di vista fusionista, nell'ambito della didattica della matematica, fu recepito da B.Rizzi e arricchito dalle sue esperienze didattiche, prima come docente nei licei e successivamente come docente universitario, nel cui ambito ha percorso tutti i gradi della carriera accademica. Dall'articolo "Interdisciplinarietà e nuovo fusionismo" [2] cito testualmente: "l'identità fra le varie discipline matematiche ed in particolare fra discipline analitiche e geometriche, si realizza per Klein non appena si pone il concetto di *gruppo di trasformazioni*. Dunque è questo il concetto primitivo, naturale, che ci consente non tanto di <<costruire>> la matematica, com'è per i bourbakisti, quanto di interpretarla" e inoltre "l'interdisciplinarietà, e in particolare il fusionismo, per quanto riguarda la genesi matematica di esso, prende atto dell'esistenza delle strutture e le usa per interpretare e presentare la matematica come un unico sistema". Ancora dallo stesso articolo: "il punto di vista del fusionista è quello in cui ogni ente è concepito non come una vuota entità formale, ma come un nome comune in cui si possono identificare, volta per volta, tutte le entità concrete delle applicazioni".

Ovviamente il Nostro non sottovaluta le difficoltà didattiche; dall'articolo "Conversazione sul Romulus" [3] cito testualmente "Purtroppo l'insegnare per problemi richiede un certo abito mentale che oggi non è molto diffuso. Bisognerebbe, in primo luogo, <<educare>> gli insegnanti ad aggredire il problema, allenarli a prospettarselo e a risolverlo sotto più punti di vista, arricchire e affinare la loro intuizione per far scoprire loro varie strade, in modo da scegliere la più semplice o la più istruttiva, mettendole poi a confronto".

Oltre gli articoli già citati, voglio ricordare, tra gli altri, i seguenti: "Idee e prassi operative, ovvero come si lavora con i concetti" [4] e "Per un nuovo ruolo della didattica" (scritto in collaborazione con il preside A.Gribaudo) [5] e ancora "La matematica tra «mondo di carta» e «mondo reale»" [6] e "L'insegnamento del calcolo delle probabilità nella scuola secondaria superiore" [7].

Infine, anche se non dedicati in modo specifico alla didattica della matematica, meritano di essere ricordati gli articoli: "Nell'educazione tecnica è basilare capire per fare" [8] e l'articolo "Chi ha detto che la tecnica è solo un fatto manuale?" [9].

Parallelamente alla pubblicazione dei suoi scritti, inesausto e ricco di contributi didattici è stato l'impegno del Nostro nella Mathesis, dagli anni sessanta fino alla sua prematura scomparsa; da socio divenne prima vicepresidente e successivamente presidente della Mathesis (1987-1993) e direttore del "Periodico di matematiche"; tale attività si esplicò nel dare i

suoi contributi didattici sotto forma di relazioni e interventi ai vari convegni provinciali e nazionali della Mathesis e di altre associazioni. Ovunque portava il contributo della sua esperienza e delle sue idee; prendendo spesso lo spunto dalla sua amata teoria dei numeri, affascinava l'uditorio con le sue argomentazioni, semplici nei contenuti strettamente matematici (perché rivolte quasi sempre a docenti di scuole primaria e secondaria), ma così profonde e ricche dal punto di vista della didattica.

Vivo è il ricordo della sua relazione al Convegno Provinciale della Mathesis di Padova del 1990. Il Nostro iniziò il suo intervento ricordando il celebre episodio citato da Hardy nel suo libro [10] dedicato a Ramanujan (entrambi celebri cultori di teoria dei numeri), e che riguardava il numero $1729 = 7 \cdot 13 \cdot 19$. Tale numero era assolutamente non significativo per Hardy, ma Ramanujan immediatamente replicò: "No, è molto interessante perché è il più piccolo numero che si può esprimere in due modi diversi come somma di due cubi di interi (infatti $1729 = 10^3 + 9^3 = 12^3 + 1^3$)". Da questo episodio, che dimostra come si possano trovare proprietà aritmetiche inaspettate per un generico numero intero, B.Rizzi, con considerazioni elementari ma didatticamente suggestive, passò a dimostrare, con semplici passaggi, il teorema che asserisce che ogni numero intero, non del tipo 2^n , $n \in \mathbb{N}_0$, si può esprimere come somma di due o più numeri interi positivi e consecutivi.

La brevità del tempo non mi consente di ricordare tanti altri episodi che rimangono come un tesoro nel mio ricordo e che, d'altra parte, sono noti ai tanti che hanno avuto la fortuna di ascoltarlo.

Negli ultimissimi anni aveva quasi cessato di frequentare i convegni e di tenere relazioni; lo assorbiva in modo totale sia il suo impegno nella lotta contro il male incurabile che aveva colpito la sua adorata Lella, sia il suo dovere di docente universitario (ha tenuto le lezioni ai suoi allievi del corso di "Metodi matematici per l'ingegneria" fino al giorno precedente la sua improvvisa scomparsa avvenuta il 6-10 u.s.).

A tutti noi, a me in particolare che ho avuto la ventura di accompagnarlo nella notte al pronto soccorso, fino all'ultimo lucido e sereno (ai miei incoraggiamenti, già sulla lettiga all'uscita del pronto soccorso, rispose con una forte stretta di mano), lascia un vivo ricordo che, citando le parole di E.Pessa, si può sintetizzare in quello di "un maestro di vita e di scienza, un uomo e un amico di straordinaria bontà e generosità, un ricercatore di eccezionale intelligenza e versatilità".

Egli sarà per noi, impegnati a vario titolo nel campo dell'insegnamento della matematica, un esempio da seguire e cercare di imitare; ciao Bruno.

BIBLIOGRAFIA

1. B.RIZZI. *de Finetti e il Periodico di Matematiche*. Periodico di Matematiche N.2/3, 1995, 69-76.
2. F. AMBRISI-B.RIZZI. *Interdisciplinarietà e nuovo fusionismo*. Archimede N.1/2, 1978, 49-51.
3. B. DE FINETTI-B.RIZZI. *Conversazione sul "Romulus"*. Periodico di Matematiche N.3/4, 1977, 21-37.
4. B.RIZZI. *Idee e prassi operativa ovvero come si lavora con i concetti*. Scuola e professionalità N.8/9, 1978, 33-37.
5. B.RIZZI-A.GRIBAUDO. *Per un nuovo ruolo della didattica*. Scuola e professionalità N.1, 1979, 28-32.
6. B.RIZZI. *La matematica tra «mondo di carta» e «mondo reale»*. Servizio informazioni Avio N.3/4, 1980, 111-116.
7. M.CESAROLI-B.RIZZI. *L'insegnamento del calcolo delle probabilità nella scuola secondaria superiore*. Atti del convegno Mathesis Paestum 18-22/4/1983, 77-85.
8. B.RIZZI-F.MORICONI. *Nell'educazione tecnica è basilare "capire" per fare*. Mensile Scuola e famiglia.
9. B.RIZZI-F.MORICONI. *Chi l'ha detto che la tecnica è solo un fatto manuale?* Mensile Scuola e famiglia.
10. G.H.HARDY. *Ramanujan*. New York-Chelsea, 1940.

RAPPRESENTAZIONE VETTORIALE DELLE SINUSOIDI

C. Di Foggia, R. Prosperì

SUNTO - Questo lavoro si propone di presentare le corrispondenze tra le sinusoidi e i vettori e tra i vettori ed i numeri complessi. Tali corrispondenze, oltre ad avere un loro interesse prettamente teorico, consentono anche notevoli semplificazioni nelle applicazioni di tali enti, soprattutto in materie tecniche quali elettrotecnica, elettronica, telecomunicazioni e sistemi.

ABSTRACT - In this work corresponding sinusoids-vectors and vectors-complexes numbers are presented. These correspondences, beyond to have a theoretical particular interest, consent also a lot of simplifications in these beings, above all in technical sciences like electro technology, electronics, telecommunications and systems.

1. Introduzione

Alcune operazioni elementari su sinusoidi, come prodotto o divisione per uno scalare o per un **operatore vettoriale**, cioè per un vettore costante nel tempo (Fig 1), vengono effettuate, in generale, in maniera alquanto semplice; altre operazioni, come la somma di sinusoidi, risultano molto complicate.

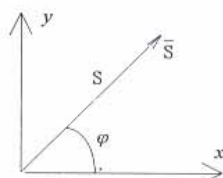


Fig. 1 - Operatore vettoriale

Consideriamo la funzione sinusoidale

$$e(t) = E \sin(\omega t + \varphi); \quad (1.1)$$

E viene detta **ampiezza** della sinusoide, ω **pulsazione**, $(\omega t + \varphi)$ **angolo di fase** o, semplicemente, **fase**. La fase è composta da due termini: ωt , variabile nel tempo, e φ , costante, coincidente con la fase all'istante iniziale $t = 0$, e, pertanto, detto **fase iniziale**; in tale istante il valore della sinusoide è $e(0) = E \sin \varphi$.

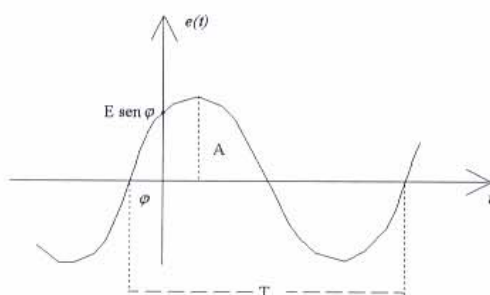


Fig. 2

Se $\varphi = 0$, la funzione (1.1) si scrive $e(t) = E \sin \omega t$ e viene presa come sinusoide di riferimento essendo, per $t = 0$, $e(0) = E \sin(0) = 0$.

In quel che segue, nel riferimento cartesiano, considereremo sull'asse delle ascisse, la variabile $\vartheta = \omega t$, invece del tempo t , in modo da leggere su tale asse direttamente la fase della sinusoide; in tal caso, quindi, la distanza tra l'origine degli assi e il punto di intersezione della sinusoide con l'asse delle ascisse rappresenta proprio la fase iniziale φ .

Rispetto alla sinusoide di riferimento, una sinusoide è:

in anticipo se $\varphi > 0$,

Fig.3

in ritardo se $\varphi < 0$.

Fig.4

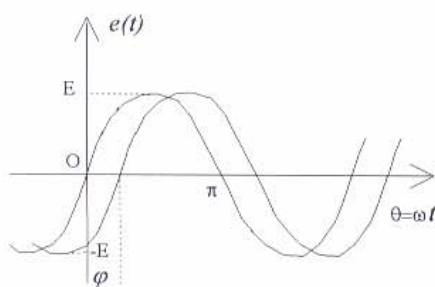


Fig.3

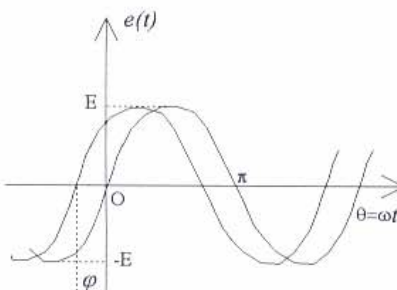


Fig. 4

2. Operazioni con sinusoidi

Riferendoci alla (1.1), con $\varphi = 0$, consideriamo alcune operazioni sulle sinusoidi:

- **moltiplicazione di uno scalare S per la sinusoide $e(t) = E \sin(\omega t)$** : dà origine ad un'altra sinusoide che è diversa dalla precedente solo per l'ampiezza

$$S \cdot e(t) = S \cdot E \sin(\omega t);$$

- **divisione della sinusoide $e(t)$ per uno scalare S** :

$$\frac{e(t)}{S} = \frac{E}{S} \sin(\omega t);$$

- **moltiplicazione della sinusoide $e(t)$ per un operatore vettoriale $\bar{S}(S, \varphi)$** : dà origine ad una sinusoide di ampiezza $S \cdot E$ e fase $\omega t + \varphi$:

$$\bar{S} \cdot e(t) = S \cdot E \sin(\omega t + \varphi).$$

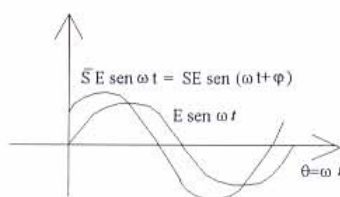


Fig. 5

- **divisione della sinusoide per l'operatore vettoriale $\bar{S}$** : la sinusoide risultante sarà

$$\frac{e(t)}{S} = \frac{E}{S} \sin(\omega t - \varphi).$$

Nel caso di **somma tra sinusoidi**, $a(t) = A \sin(\omega t)$ e $b(t) = B \sin(\omega t + \varphi)$, occorrerebbe eseguire i seguenti laboriosi calcoli:

$$a(t) + b(t) = A \sin(\omega t) + B \sin(\omega t + \varphi) = A \sin \omega t + B [\sin \omega t \cos \varphi + \cos \omega t \sin \varphi] =$$

$$= (A + B \cos \varphi) \left[\sin \omega t + \frac{B \sin \varphi}{A + B \cos \varphi} \cos \omega t \right]; \text{ ponendo } \frac{B \sin \varphi}{A + B \cos \varphi} = \tan \alpha :$$

$$(A + B \cos \varphi) \left[\sin \omega t + \frac{\sin \alpha}{\cos \alpha} \cos \omega t \right] = \frac{(A + B \cos \varphi)}{\cos \alpha} [\sin \omega t \cos \alpha + \sin \alpha \cos \omega t] =$$

$$= K \sin(\omega t + \alpha), \text{ avendo posto } K = \frac{A + \cos \varphi \cdot B}{\cos \alpha}.$$

La sinusoide somma ha la stessa pulsazione delle componenti ma fase e ampiezza diverse, quindi la somma è ancora una sinusoide; infatti, essa è, tra le funzioni periodiche, l'unica che conserva la propria forma.

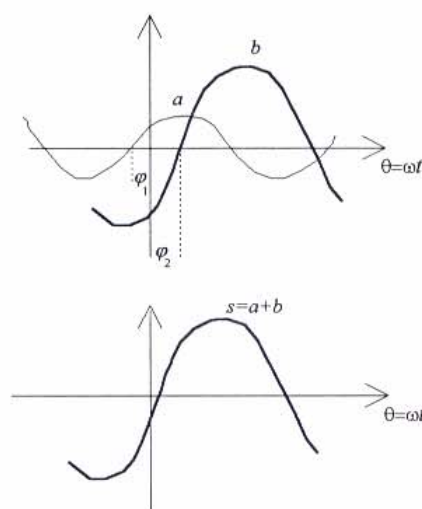


Fig. 6

3. Le sinusoidi come vettori

Supponiamo di rappresentare i vettori in un riferimento cartesiano Oxy ; consideriamo un **vettore** $\vec{V}$ **rotante** con velocità angolare ω costante, modulo V e fase iniziale φ ; se all'istante $t=0$ ne proiettiamo l'estremo sull'asse delle y , otteniamo su di esso un segmento di lunghezza $V \sin(\varphi)$; su un piano (t,y) otteniamo il punto $(0, V \sin(\varphi))$. Dopo un tempo t_1 il vettore ruota di un angolo di fase pari a ωt_1 ; la sua fase complessiva è pari a $\omega t_1 + \varphi$ e la sua proiezione sull'asse y è $V \sin(\omega t_1 + \varphi)$; il corrispondente punto nel piano (t,y) è $(t_1, V \sin(\omega t_1 + \varphi))$. Tali punti hanno un andamento sinusoidale nel tempo, con periodo $T = \frac{2\pi}{\omega}$.

Dopo un tempo $t_1 = \frac{\vartheta_1}{\omega}$ il vettore ha ruotato di un angolo ϑ_1 , per cui la sua fase è $(\vartheta_1 + \varphi)$; la sua proiezione sull'asse y è

$$V \sin(\omega t_1 + \varphi) = V \sin(\vartheta_1 + \varphi)$$

e l'ascissa corrispondente è $t_1 = \frac{\vartheta_1}{\omega}$ o $\vartheta_1 = \omega t_1$.

Quindi, il punto sul piano (ϑ, y) è $(\vartheta_1, V \sin(\vartheta_1 + \varphi))$.

L'ampiezza della sinusoide è uguale al modulo del vettore e la sua fase iniziale è uguale all'angolo, computato in senso antiorario, che il vettore forma con l'asse di riferimento x nell'istante zero; questi due elementi sono sufficienti ad identificare la sinusoide.

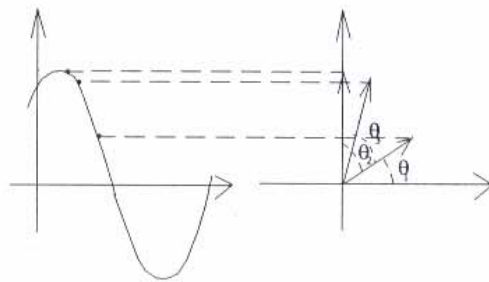


Fig 7

Ovviamente, ad una sinusoide avente una certa ampiezza ed una certa fase possiamo far corrispondere un vettore con modulo pari all'ampiezza della sinusoide e ruotato di un angolo φ rispetto all'asse delle x .

È ovvio che solo i vettori rotanti rappresentano delle sinusoidi, i non rotanti sono gli **operatori vettoriali**.

Nella rappresentazione delle sinusoidi abbiamo utilizzato sull'asse di riferimento $\vartheta = \omega t$ in modo tale da poter leggere direttamente il valore di φ . Se considerassimo come asse di riferimento quello dei tempi dovremmo di volta in volta calcolare i prodotti

$$\vartheta_1 = \omega t_1, \quad \vartheta_2 = \omega t_2, \quad \vartheta_3 = \omega t_3.$$

Considerare al posto della sinusoide la sua rappresentazione vettoriale ci facilita i calcoli da eseguire; sia $\vec{V}$ il vettore che rappresenta la sinusoide (1.1)

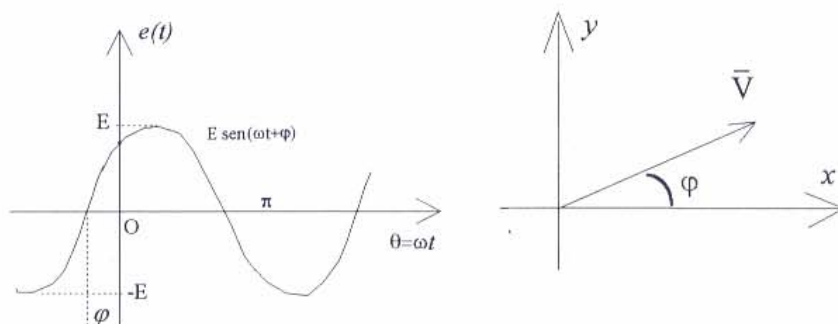


Fig. 8

- **Moltiplicazione di uno scalare S per il vettore $\bar{V}$ che rappresenta la sinusoide.** Il prodotto di uno scalare per un vettore dà origine ad un vettore di modulo SV e fase uguale a quella di $\bar{V}$.
- **Divisione del vettore $\bar{V}$ per uno scalare.** Il rapporto fornisce un vettore di modulo pari al rapporto dei moduli e fase pari alla differenza delle fasi.
- **Moltiplicazione di un operatore vettoriale $\bar{S}$ per il vettore sinusoide $\bar{V}$.** Il modulo del loro prodotto è dato dal prodotto dei moduli e la fase è uguale alla somma delle fasi.

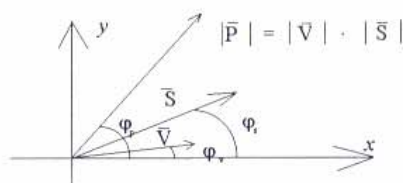


Fig. 9

- **Divisione di un operatore vettoriale $\bar{S}$ per il vettore sinusoide $\bar{V}$.**
- **Somma fra due sinusoidi utilizzando i rispettivi vettori**

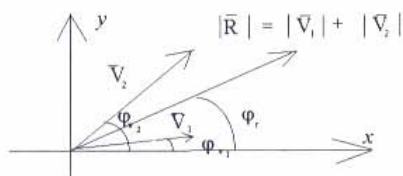


Fig. 10

Nella costruzione geometrica che dà origine alla sinusoide, si considera un punto P sulla circonferenza di raggio unitario $\overline{OP}$.

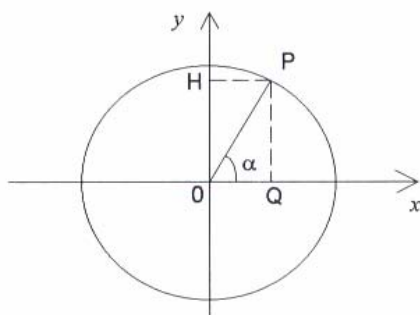


Fig. 11

Sia α l'angolo percorso in senso antiorario che il segmento OP forma con l'asse delle ascisse, sia OH la proiezione di OP sull'asse delle y. Riportando, al variare di α , queste proiezioni su di un sistema di assi cartesiani otteniamo la sinusoide.

Il discorso è analogo se anziché considerare un punto che si muove (che descrive una circonferenza) si sostituisce ad esso un vettore OP, di modulo uguale ad 1, rotante con velocità angolare ω intorno ad un asse perpendicolare al piano e passante per O.

4. Vettori e numeri complessi

Oltre che dal punto di vista grafico, i vettori possono essere trattati analiticamente utilizzando la corrispondenza con i numeri complessi. È noto che un numero complesso $z = a + jb$ individua il vettore del piano $\vec{z}(a,b)$ di componenti rispettivamente a e b , dove a è la parte reale, b la parte immaginaria e $j = \sqrt{-1}$ è l'unità immaginaria che nel piano cartesiano rappresenta il versore dell'asse delle ordinate $\vec{j}(0,1)$ che individua (nella moltiplicazione tra numeri complessi) la rotazione di un angolo uguale a $\frac{\pi}{2}$ radianti.

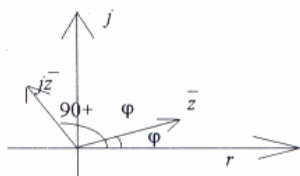


Fig. 12

Il piano dei numeri complessi è detto *Piano di Argand-Gauss*; in esso l'asse delle ascisse prende il nome di *retta dei numeri reali* e quello delle ordinate prende il nome di *retta dei numeri immaginari*.

Ad ogni numero complesso corrisponde, in tale piano, un punto avente per ascissa la propria parte reale e per ordinata il proprio coefficiente dell'immaginario. Quindi, a ciascun numero complesso $\bar{z} = a + jb$ si associa un vettore, indicato con il simbolo $\bar{z}$, che unisce l'origine al punto di coordinate a e b e che avrà modulo $z = |\bar{z}| = \sqrt{a^2 + b^2}$ e fase $\varphi = \arg(\bar{z}) = \arctg \frac{b}{a}$.

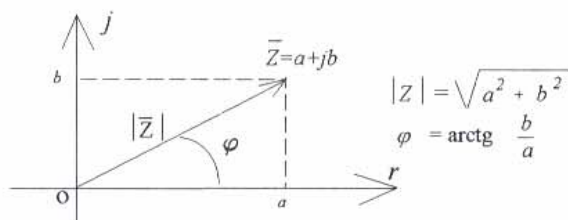


Fig. 13

Con le operazioni tra vettori si possono eseguire le operazioni di somma e differenza fra grandezze alternate.

Ad esempio, la somma di due correnti in un circuito elettrico in regime sinusoidale permanente si può ottenere nel modo seguente: si associa ad ognuna delle due sinusoidi il vettore corrispondente, si ricava il risultante dei vettori e si associa ad esso la sinusoide corrispondente.

La rappresentazione vettoriale consente di evidenziare in maniera chiara e immediata (cioè senza tracciare i grafici) le relazioni esistenti tra le fasi di funzioni (segnali) sinusoidali.

Un esempio di ciò può essere costituito dal calcolo della differenza di fase esistente tra corrente e tensione di un condensatore.

BIBLIOGRAFIA

1. G. BOBBIO e S. SAMMARCO, *Elettrotecnica generale con esercitazioni di laboratorio e misure*, PETRINI, Torino, 1989.
2. G. MARTINELLI e R. SALERNO, *Fondamenti di elettrotecnica vol.I*, Edizioni Scientifiche Siderea, Roma, 1979.

CAMPI DI GALOIS: UNA PRESENTAZIONE DIVULGATIVA

Fernando Di Gennaro^{*}

Il presente lavoro è dedicato alla memoria del mio Maestro prof. Bruno Rizzi, immaturamente scomparso. Il suo ricordo sarà con me per sempre.

PREMESSA

Per molto tempo il nome di Galois, come la sua opera non è entrato negli argomenti di insegnamento. Negli ultimi trent'anni è nato un notevole interesse volto a riscoprire, ad ampliare e ad approfondire alcuni aspetti delle sue ricerche. Parimenti nasce, di conseguenza, il problema della divulgazione delle parti più semplici della sua opera...

Questo articolo propone in sostanza una scelta di argomenti ed una metodologia, molto operativa, fruibile fin da una terza Liceo Scientifico, atta a permettere agli allievi di impadronirsi dei concetti basilari della Teoria dei Campi Finiti, o "Campi di Galois", e di conseguenza delle Geometrie finite che su di essi poggiano.

^{*} Liceo Scientifico *M. Curie* di Giulianova (TE)

1. LA NOZIONE DI CAMPO

Sia Q un insieme contenente almeno due elementi distinti che formalmente indichiamo con i simboli 0 e 1 . Supponiamo che in tale insieme Q siano definite due operazioni binarie interne, dette rispettivamente la prima "addizione" e denotata con il simbolo formale $(+)$, la seconda "moltiplicazione", denotata simbolicamente con $(\bullet)$.

La struttura $Q = (Q, +, \bullet)$ si dice che è un *corpo* se verifica i seguenti assiomi:

- 1.- La struttura $(Q, +)$ è un gruppo abeliano, cioè $\forall a, b, c \in Q$ valgono le seguenti proprietà:

I	$(a + b) + c = a + (b + c)$	[proprietà associativa]
II	$a + b = b + a$	[proprietà commutativa]
III	$\exists 0 \in Q, \text{ t.c. } 0 + a = a + 0 = a$	[esistenza elemento neutro]
IV	$\forall a \in Q, \exists -a \in Q \text{ tale che: } a + (-a) = (-a) + a = 0$	[esistenza dell'opposto]

- 2.- La struttura $(Q \setminus \{0\}, \bullet)$ è pure un gruppo, ma non necessariamente commutativo.

Cioè $\forall a, b, c \in Q$, valgono le seguenti proprietà:

I	$(a \bullet b) \bullet c = a \bullet (b \bullet c)$	[proprietà associativa]
II	$\exists 1 \in Q, \text{ t.c. } 1 \bullet a = a \bullet 1 = a$	[esistenza neutro moltiplicativo]
III	$\forall a \in Q, \exists a' \in Q \text{ tale che: } a \bullet a' = a' \bullet a = 1$	[esistenza elemento inverso]

- 3.- Infine la moltiplicazione è distributiva rispetto all'addizione, cioè $\forall a, b, c \in Q$,

$$a \bullet (b + c) = a \bullet b + a \bullet c.$$

Quando la moltiplicazione è commutativa il corpo si dice che è un *campo*, altrimenti si dice un *corpo sghembo*. Sono esempi di campi l'insieme dei numeri razionali, dei numeri reali, dei numeri complessi. Fra i campi ci sono quelli infiniti come i campi numerici sopracitati; ci sono anche quelli finiti che sono l'oggetto della presente nota.

Si chiama *campo di Galois* un campo che contiene un numero finito q di elementi. Esso si suole indicare con $GF(q)$. Si può dimostrare il seguente notevole:

TEOREMA DI VEDDEBURN. *Ogni corpo finito è un campo.*

Il numero q , che è il numero degli elementi del campo, si chiama *ordine del campo di Galois*. Si possono dimostrare, cosa che noi non faremo, i

seguenti tre teoremi che caratterizzano in un certo qual senso i campi di Galois. Il teorema che segue fornisce tutte le condizioni, in semplici termini di cardinalità, relative alla esistenza di un campo di Galois.

TEOREMA DI ESISTENZA ED UNICITA' DEI CAMPI FINITI. *Un campo di Galois K (finito) d'ordine q esiste se e solo se l'intero q è una potenza di un numero primo p , cioè*

$$q = p^h$$

Inoltre due campi finiti di eguale ordine sono necessariamente isomorfi.

Il teorema sopradetto ci dice tutto tranne come costruire operativamente tali campi finiti. Questa parte costruttiva è quella che ci interessa maggiormente per accostare gli allievi alla teoria. Le prove dei due teoremi sopra enunciati sono per così dire di livello superiore alle unità didattiche che stiamo proponendo. Noi ci occuperemo in questo lavoro di costruire in un primo momento campi finiti di ordine primo, ovvero di introdurre l'allievo alla aritmetica modulare, in una fase successiva del lavoro ci occuperemo effettivamente della costruzione dei campi di Galois aventi come ordine una potenza di un primo. Per questi ci basterà sapere che fissato l'ordine essi sono unici, a meno di isomorfismi.

2. CAMPI DI ORDINE PRIMO

Al fine di costruire i campi aventi come ordine un numero primo, è necessario introdurre l'aritmetica modulare. Sia $\mathbf{Z}$ l'insieme dei numeri interi relativi e siano n e h due interi relativi verificanti le condizioni

$$n \in \mathbf{Z}, n > 2 \text{ e } 0 \leq h \leq n-1.$$

Si indichi con $[h]$ l'insieme dei numeri interi relativi che divisi per n danno come resto h , cioè l'insieme dei numeri del tipo $kn + h$. A titolo di esempio se prendiamo $n = 6$ allora la classe $[3]$ è costituita da tutti i numeri della forma $6 \cdot k + 3$ cioè dai numeri che divisi per 6 danno resto 3, dunque i numeri 3, 9, 21, 27, ad esempio, sono tutti nella classe $[3]$ (modulo 6), essendo:

$$0 \cdot 6 + 3 = 3, 1 \cdot 6 + 3 = 9, 2 \cdot 6 + 3 = 15, 3 \cdot 6 + 3 = 21, 4 \cdot 6 + 3 = 27, \dots$$

Si ottengono, fissato n , esattamente n classi modulo n e precisamente:

$$[0], [1], [2], \dots, [n-1].$$

Tali classi sono a due a due disgiunte ed ogni numero intero relativo appartiene ad una e una sola di tali classi. Le classi formano allora quella che si chiama una partizione di $\mathbf{Z}$. Ciascuna classe si chiama una classe resto modulo n o un intero ridotto modulo n . L'insieme delle classi modulo n si denota con $\mathbf{Z}(n)$.

Se $n = 2$ l'insieme $\mathbf{Z}(2)$ ha come elementi la classe $[0]$ che è costituita da tutti i numeri pari e la classe $[1]$ che è costituita da tutti i numeri dispari.

Nell'insieme $\mathbf{Z}(2)$ si possono definire, $\forall [a], [b] \in \mathbf{Z}(2)$ due operazioni nel modo che segue:

$$[a] + [b] := [a + b]; \quad [a] \bullet [b] := [a \bullet b]$$

La definizione nuova di addizione e moltiplicazione tra classi che appare al primo membro è definita mediante l'addizione e la moltiplicazione di elementi di $\mathbf{Z}$. Sarebbe bene a questo punto adoperare del tempo con gli allievi sia per esemplificare, sia per mostrare loro come le operazioni siano indipendenti dai rappresentanti scelti, principalmente capire bene questa idea. Sempre nella fase di esercitazioni è poi bene dedicare il giusto tempo a verificare che quale che sia n la struttura $\mathbf{Z}(n)$ ha come zero ed uno sempre le classi $[0]$ ed $[1]$.

Il nostro interesse successivo è l'esame più completo della struttura $\mathbf{Z}(n)$. E' piuttosto interessante dilungarsi sulle prove dei fatti seguenti:

- A) Verifica del fatto che $(\mathbf{Z}(n), +)$ è un gruppo commutativo, verificandone le varie proprietà.
- B) Verifica per la struttura $(\mathbf{Z}(n), \bullet)$ della proprietà associativa e commutativa e verifica della distributiva.

Vale poi la pena dedicare tempo alla dimostrazione del seguente

TEOREMA. *La struttura $\mathbf{Z}(n)$ possiede divisori dello zero se e solo se n è composto.*

DIMOSTRAZIONE. Sia n composto, esistono allora due divisori propri di n . Dunque si ha $n = a \bullet b$ con $1 < a, b < n$. Segue

$$[a] \bullet 1[b] = [n] = [0] \quad \text{con } [a], [b] \neq [0]$$

Cioè se n è composto la struttura ha divisori dello zero.

Inversamente, se risulta $[a] \bullet 1[b] = [0]$ con $[a], [b] \neq [0]$ allora risulta $[a \bullet 1b] = [n]$ con $1 < a, b < n$ e quindi $n = a \bullet b$ con $1 < a, b < n$ e quindi n è composto.

Quanto provato si completa con il seguente

TEOREMA. *Un elemento di $\mathbf{Z}_n$ è invertibile se e solo se non è un divisore dello zero.*

DIMOSTRAZIONE. Supponiamo che in un qualsiasi anello l'elemento $a \neq 0$ sia un divisore proprio dello zero (allora sia $b \neq 0$ un divisore complementare) e che l'anello abbia almeno due elementi.

Se allora si suppone che a abbia anche inverso a' segue:

$$\begin{aligned} a'ab &= (a'a)b = 1b = b \\ a'ab &= a'(ab) = a'0 = 0 \end{aligned}$$

da cui $b = 0$ contro l'ipotesi. Dunque un divisore dello zero non può essere invertibile ed un elemento invertibile non può esser divisore dello zero essendo le due relazioni contraddittorie.

Rimane da provare che ogni elemento di Z che non è un divisore dello zero è effettivamente invertibile.

Allo scopo sia a un intero che non divide n (cioè che non sia un divisore dello zero) e tale che $1 < a < n$. Considero i multipli ta di a , per $t = 1, 2, \dots, n$; vogliamo provare che uno di essi è divisibile per n . Dividiamo con resto il generico ta per n , si ha:

$$ta = rn + r$$

Supponiamo $t = 1, 2, \dots, n$, siano t ed s due valori tra questi, risulta allora $r = r$ se e solo se $t = s$ altrimenti si avrebbe:

$$(t-s)a = (r-r)n + (r-r) = (r-r)n$$

con $|(t-s)a| < n$ e $|(r-r)n| > n$.

Dunque gli n resti sono tutti distinti ed uno di essi è quindi nullo.

Questa breve e semplice prova può essere fortemente istruttiva e permette per via elementare di provare l'esistenza dell'inverso. Si potrebbe anche tentare di provare ai ragazzi qualche cosa in più come ad esempio che ogni anello finito privo di divisori dello zero è un corpo, ma questo è da sperimentare. A questo punto occorre lavorare con le tabelle e impadronirsi delle classi resto in modo realmente operativo.

La conclusione comunque del nostro paragrafo è riassunta dal seguente

TEOREMA. *La struttura Z_n avente come elementi le classi resto, munita delle operazioni di addizione e moltiplicazione, è un campo se e solo se n è primo.*

Nel seguito porremo per ogni primo p : $GF(p) = Z_p$. La costruzione dei campi di Galois di ordine primo è così completamente assegnata.

3. CAMPI DI ORDINE POTENZA DI UN PRIMO

Partiamo con dei crudi esempi costruttivi.

Si costruisca $GF(9)$. Partiamo da $Z_3 = Z(3)$ formato dalle classi 0, 1, 2 (nel seguito ometteremo le parentesi quadra per semplicità) con

$$\begin{aligned} 0 + 0 &= 2 + 2 = [0], & 1 + 0 &= 2 + 2 = [1], & 1 + 1 &= 2 + 0 = [2] \\ 1 * 2 &= 2 & 2 * 2 &= 1 * 1 = 1 \end{aligned}$$

Consideriamo l'equazione

$$x^2 + 1 = 0$$

Tale equazione in $Z(3)$ è priva di soluzioni poiché i tre numeri

$$\begin{aligned} 0 * 0 + 1 &= 1 \\ 1 * 1 + 1 &= 2 \\ 2 * 2 + 1 &= 2 \end{aligned}$$

non sono mai nulli.

Si considerino i binomi del tipo $a + ib$ dove a, b variano in $Z(3)$ ed i è una "soluzione formale" dell'equazione fissata.

Operando poi sui binomi con la somma e il prodotto, si ha:

$$\begin{aligned} (2 + i) + (2 + 2i) &= (1 + i) * (1 + 2i) = 1 + 4 = 2 \\ i * (1 + i) &= 2 + i \end{aligned}$$

L'insieme di questi binomi con le operazioni è una descrizione di $GF(9)$. Cambiando l'equazione cambia la descrizione, ma si ha sempre $GF(9)$ a meno di isomorfismi.

Data l'equazione in $Z(2)$

$$x^2 + x + 1 = 0$$

possiamo costruire gli elementi di $GF(4)$ dati da: 0, 1, i , $1+i$.

Per essi si ha $i = i + 1$.

Con l'equazione

$$x^2 + 2x + 1 = 0$$

si può costruire $GF(27)$ i cui elementi sono trinomi del tipo $ai + bi + c$ con a, b, c e $Z(3)$ e dove la somma dei trinomi è quella usuale con la sola sostituzione $i = -2i - 1 = i + 2$.

Con l'equazione

$$x^2 = x + 1$$

si costruisce $GF(16)$ che ha come elementi dei polinomi di 3° ordine.

4. LA DIVISIONE IN UN CAMPO DI GALOIS

Il punto di partenza ovvio, ma da rimarcare con i ragazzi, è far capire che dividere è moltiplicare un elemento per il suo inverso. Questo ha senso essenzialmente per il fatto che la struttura moltiplicativa è commutativa. In altri ambienti come quello dell'algebra delle matrici non ha senso fare un quoziente, ma si può fare un quoziente destro ed uno sinistro...

Ad esempio, a che cosa deve corrispondere 1 diviso 2, cioè $1/2$, in questo campo finito? Deve essere necessariamente uno dei numeri 0, 1, 2, 3, 4. Ma quale? In altre parole, qual è quel numero che moltiplicato per 2 dà 1?

Possiamo procedere per tentativi:

$$\begin{aligned} 2 \times 0 &= 0 \neq 1 \\ 2 \times 1 &= 2 \neq 1 \\ 2 \times 2 &= 4 \neq 1 \\ 2 \times 3 &= 6 = 1 \end{aligned}$$

Allora possiamo identificare $1/2$ con 3. Però non è possibile procedere sempre per tentativi per trovare l'inverso di un numero: c'è una legge generale che in questo caso è:

$$1/a = a^3$$

Infatti

$$\begin{aligned} 1/2 &= 2^3 = 8 = 3 \\ 1/4 &= 4^3 = 64 = 4 \end{aligned}$$

Quindi 4 è il reciproco di se stesso.

Più in generale, per trovare il reciproco di un numero si può considerare la congruenza

$$a^{\Phi(m)} \equiv 1 \pmod{m}$$

che possiamo scrivere nella forma

$$a \bullet a^{\Phi(m)-1} = 1$$

da cui

$$1/a = a^{\Phi(m)-1}$$

dove $\Phi(m)$ è l'indicatore di Eulero.

I campi finiti si sono assicurati un posto in algebra per merito di Galois che li usò per mostrare le condizioni sotto cui le equazioni algebriche hanno soluzioni per radicali.

5. L'EQUAZIONE DI SECONDO GRADO IN UN CAMPO DI GALOIS

Ci interessiamo ora della risoluzione delle equazioni di 2° grado in un campo di Galois $GF(q)$. Per semplicità d'ora in poi scriveremo gli elementi di $GF(q)$ omettendo le parentesi quadre. Quindi a al posto di $[a]$.

Sia ora

$$x^2 = a \quad (1)$$

Proviamo che un'equazione di 2° grado pura ha al più due soluzioni tra loro opposte.

Iniziamo con il notare che la (1) non ha soluzioni quando a non è un quadrato. Ha inoltre una sola soluzione se $a = 0$, come è ovvio (altrimenti ci sarebbero divisori dello zero).

Supponiamo a quadrato, cioè supponiamo che esista $b \in GF(q)$ tale che $a = b^2$. Allora la (1) ha almeno due soluzioni $x_1 = b$, $x_2 = -b$ tra loro opposte.

Proviamo che non ce ne sono altre.

Sia per questo x_3 una eventuale ulteriore soluzione con $x_3 \neq x_1$ ed $x_3 \neq x_2 = -x_1$. Risulta cioè $x_3 - x_1 \neq 0$ ed $x_3 + x_1 \neq 0$. Avendosi allora

$$x_1^2 + x_3^2 = \Delta$$

segue:

$$0 = x_3^2 - x_1^2 = (x_3 - x_1)(x_3 + x_1)$$

Cioè l'esistenza di una terza radice implicherebbe l'esistenza di divisori dello zero.²

Siamo ora in grado di trattare il caso generale:

$$ax^2 + bx + c = 0.$$

Si può scrivere

$$ax^2 + bx + c = a \left[\left(x + \frac{b}{2a} \right)^2 - \frac{\Delta}{4a^2} \right]$$

essendo $\Delta = b^2 - 4ac$.

Dunque il problema è ricondotto alla equazione:

$$\left(x + \frac{b}{2a}\right)^2 = \frac{\Delta}{4a^2}$$

Denotiamo, nell'ipotesi che Δ sia un quadrato, con:

$$\sqrt{\Delta} = \{r : r \in GF(q), r^2 = \Delta\}.$$

Segue allora la "formula risolutiva":

$$x = \frac{-b + \sqrt{\Delta}}{2a}$$

avendo $\sqrt{\Delta}$ due determinazioni.

Quando $\Delta = 0$ l'equazione ha la "radice doppia" espressa sempre dalla formula, mentre nel caso " Δ non quadrato" non ci sono soluzioni.

Il procedimento indicato cade in difetto per il caso di q pari.

Supponiamo ora che q sia pari. Si prova subito che $\forall a \in GF(q)$ è:

$$a^2 = b^2 \Leftrightarrow a = b$$

Infatti da $a^2 - b^2 = 0$ risulta $a^2 + b^2 = 0$ per essere $x = -x$. Inoltre essendo $2 = 0$ e quindi $2ab = 0$ è:

$$a^2 + b^2 = a^2 - 2ab + b^2 = (a - b)^2 = 0$$

da cui l'asserto.

Segue allora che i quadrati degli elementi di $GF(q)$ sono a due a due distinti e quindi

ogni elemento di $GF(q)$ è un quadrato.

Segue allora che una equazione della forma

$$ax^2 + bx + c = 0.$$

con $b = 0$ ha sempre una ed una sola soluzione.

Proviamo ora che una equazione del tipo

$$ax^2 + bx + c = 0.$$

con $b \neq 0$ equivale ad una equazione del tipo $y^2 + y = d$.

Infatti, posto $x = hy$ l'equazione diventa:

$$ah^2y^2 + bhy + c = 0.$$

con la condizione che sia $ah^2 = bh$ allora $ah = b$. Dividendo tutti i termini per ah si ottiene appunto l'equazione richiesta.

Si dimostra che gli elementi di $GF(q)$, q pari, si dividono in due categorie. Gli elementi di I categoria sono gli elementi d che si possono scrivere nella forma:

$$d = y^2 + y$$

Per questi l'equazione associata ha due soluzioni che sono y ed $y + 1$. Risulta infatti:

$$(y + 1)^2 + (y + 1) = y^2 + 1 + y + 1 = y^2 + y = d$$

Gli altri elementi detti di II categoria sono quelli che non ammettono la suddetta decomposizione. Per questi l'equazione associata non ha soluzione.

BIBLIOGRAFIA

1. M. CERASOLI, F. EUGENI e M. PROTASI, *Elementi di Matematica Discreta*, Zanichelli, Bologna, 1988.

NOTE SU DUE PROBLEMI DI GEOMETRIA E SUCCESSIONI RICORSIVE

Franco Mancinelli *

Dedico questo lavoro alla memoria del Prof. Bruno Rizzi.

SUNTO. La presente nota si divide in quattro paragrafi. Nel primo si considera la successione dei poligoni regolari di fissato perimetro e tali che il successivo ha numero di lati doppio del precedente; apotema e raggio di tali poligoni sono legati da formule ricorrenti il cui limite dipende esclusivamente dalla figura iniziale che si considera. Nei paragrafi 2 e 3 si studiano le ricorrenze indipendentemente dai poligoni e se ne danno nuove e più ampie interpretazioni geometriche. Nell'ultimo paragrafo si esaminano dapprima le analoghe relazioni intercorrenti tra apotema e raggio di poligoni regolari di area assegnata, poi si studiano le analoghe successioni svincolate dai poligoni.

1. PRIMO PROBLEMA

Siano a_0 e r_0 apotema e raggio rispettivamente di un poligono regolare di numero di lati n_0 e perciò di lato l_0 :

$$l_0 = \sqrt{r_0^2 - a_0^2} \quad \text{con } a_0 < r_0$$

e perimetro p_0 :

* Liceo Scientifico *M. Curie* di Giulianova (Te).

$$p_0 = 2n_0 \sqrt{r_0^2 - a_0^2}$$

E' noto ([1], pag. 161) che , considerata la successione dei poligoni isoperimetri, in cui il poligono regolare successivo ha numero di lati doppio del poligono regolare precedente , allora apotema e raggio dei suddetti poligoni regolari, a partire da $n_0=3$ sono, alternativamente , medio aritmetico e medio geometrico di apotema e raggio immediatamente precedenti ; con ovvio significato dei simboli , si ha :

$$1) \quad a_{2n} = \frac{a_n + r_n}{2} \quad r_{2n} = \sqrt{a_{2n} \cdot r_n}$$

Si pongono, per la successione 1) due problemi: il primo riguarda la sua natura ; il secondo, nel caso in cui essa sia convergente, il calcolo del limite. A proposito della prima questione, si dimostra che la successione 1) è convergente; tale risultato è , d'altra parte intuitivo , se si pensa che al tendere di n a $+\infty$, le sottosuccessioni a_n e r_n tendono entrambe al raggio l della circonferenza limite γ_l cui tende la successione dei poligoni isoperimetri . Per quanto riguarda il valore del limite l della successione 1) , esso dipende , naturalmente, dai valori dell' apotema a_0 e del raggio r_0 del poligono iniziale di numero di lati n_0 ed è calcolabile ragionando nei termini seguenti: se indichiamo con p_0 il perimetro del poligono regolare iniziale , si ha:

$$p_0 = 2n_0 \sqrt{r_0^2 - a_0^2} \quad ;$$

D'altra parte, il lato l_{2n} del poligono regolare con numero di numero di lati $2n$ è dato da :

$$l_{2n} = 2n \sqrt{r_{2n}^2 - a_{2n}^2}$$

e il perimetro p_{2n} da:

$$2) \quad p_{2n} = 4n \sqrt{r_{2n}^2 - a_{2n}^2}$$

Se con α_{2n} indichiamo la misura, in radianti, dell'angolo che

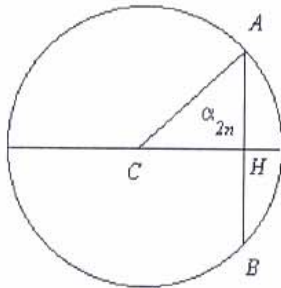


fig. 1

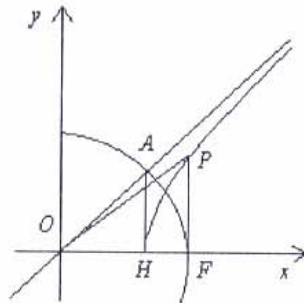


fig. 2

il raggio AC forma con l'ipotenusa CH (fig 1), si ha: $n = \pi / (2\alpha_{2n})$, e perciò dalla 2) segue :

$$p_{2n} = \frac{2\pi}{\alpha_{2n}} \sqrt{r_{2n}^2 - a_{2n}^2}$$

Poiché $\cos(\alpha_{2n}) = (r_{2n}/a_{2n})$, risulta $\alpha_{2n} = \arccos(r_{2n}/a_{2n})$, e dalle 2) e 3) seguono, nell'ordine :

$$4) \quad p_0 = \frac{2\pi}{\alpha_{n_0}} \sqrt{r_{n_0}^2 - a_{n_0}^2} \quad ,$$

$$5) \quad p_{2n} = \frac{2\pi}{\alpha_{2n}} \sqrt{r_{2n}^2 - a_{2n}^2} \quad .$$

Poiché i poligoni sono isoperimetri, non solo si ha che $p_{2n} = p_0$, ma si ha anche, variando n e perciò a_{2n} , r_{2n} secondo la 1), da una parte che l'espressione a secondo membro della 5) resta costante, dall'altra che il valore di tale espressione è eguale alla lunghezza $2\pi l$ della circonferenza limite γ_l . In simboli :

$$6) \quad \frac{2\pi \sqrt{r_0^2 - a_0^2}}{\arccos(r_0/a_0)} = \frac{2\pi \sqrt{r_{2n}^2 - a_{2n}^2}}{\arccos(r_{2n}/a_{2n})} = 2\pi l \quad ,$$

da cui :

$$7) \quad l = \frac{\sqrt{r_0 - a_0}}{\arccos(a_0/r_0)} ,$$

risultando , tale valore di l , il limite della successione ricorsiva data.

E' interessante notare che l' espressione k_{2n} così definita :

$$8) \quad k_{2n} = \frac{\sqrt{r_{2n}^2 - a_{2n}^2}}{\arccos(a_{2n}/r_{2n})}$$

rimane costante quando a_{2n} e r_{2n} variano secondo la legge 1), e rappresenta il rapporto il perimetro costante e 2π .

2. Generalizzazioni

Anche nel caso di a_0 , b_0 reali positivi , con $a_0 < b_0$, si dimostra che ([2],pag. 81) la successione ricorsiva :

$$1*) \quad a_{n+1} = \frac{a_n + b_n}{2} , \quad b_{n+1} = \sqrt{a_{n+1}^2 - b_n^2} ,$$

è convergente ; in particolare la quantità k_n , definita ponendo:

$$8*) \quad k_n = \frac{\sqrt{b_n^2 - a_n^2}}{\arccos(a_n/b_n)} ,$$

(cfr. 8)) resta costante al variare di n e il limite l della successione 1*) è dato dall' espressione

$$7*) \quad l = \frac{\sqrt{b_0^2 - a_0^2}}{\arccos(a_0/b_0)}$$

analoga alla 7).

Nel caso di $a_0, b_0 \in \mathbb{R}_0^+$ però con $a_0 > b_0$ la successione ricorsiva definita dalla 1*) è ancora convergente al limite l dato da :

$$7**) \quad l = \frac{\sqrt{a_0^2 - b_0^2}}{\arccos h(a_0/b_0)}$$

formalmente simile alla 7*), differendo da essa per il fatto che a denominatore è presente di (a_0/b_0) il coseno iperbolico e non il coseno. In tal caso, inoltre, al variare di n , rimane costante la quantità k_n :

$$8^{**}) \quad k_n = \frac{\sqrt{a_n^2 - b_n^2}}{\operatorname{arccos} h(a_n/b_n)}$$

anch' essa formalmente analoga alla omonima della 8*).

E' interessante notare che la successione ricorsiva 1*) e, in particolare, la quantità k_n data in 8**), anche in questo caso $a_0 > b_0$, sono suscettibili della seguente interessante interpretazione geometrica: a partire dall' iperbole γ_0 di semidistanza focale a_0 e semiasse trasverso b_0 , consideriamo la successione di iperboli γ_n di eq.:

$$\gamma_n: \quad \frac{x^2}{b_n^2} - \frac{y^2}{a_n^2 - b_n^2} = 1$$

in cui le quantità a_n e b_n rappresentano, ora, rispettivamente le semidistanze focali e i semiassi trasversi di γ_n ; un ramo di γ_n è riportato in fig 2; in essa è: $H(b_n; 0)$, $F(a_n; 0)$, $OA=OF=a_n$, $P(x_P; y_P) \in \gamma_n$, $x_P=a_n$, $y_P=(a_n^2 - b_n^2)/b_n = PF$; al variare di n e perciò di a_n e b_n secondo la 1*), varia la generica iperbole; resta costante la quantità k_n la cui espressione data nella 8**), rappresenta, ora, il prodotto di AH per il rapporto tra il triangolo AOH e il triangolo mistilineo OPH ; in simboli:

$$k_n = AH \frac{AOH}{OPH}.$$

Infatti è (fig. 2)

$$AH = \sqrt{a_n^2 - b_n^2},$$

mentre

$$(OPH/AOH) = \operatorname{arccos} h(a_n/b_n);$$

quest' ultima affermazione si prova facilmente se si nota, da una parte, che

$$AOH = b_n \sqrt{a_n^2 - b_n^2} / 2,$$

dall' altra, che $OPH=OPF-HPF$, con

$$OPF = a_n(a_n^2 - b_n^2) / (2b_n)$$

e con

$$HPF = \sqrt{\frac{a_n^2 - b_n^2}{b_n^2}} \int_{bn}^{an} \sqrt{x^2 - b_n^2} dx,$$

si che

$$OPH = (1/2)b_n\sqrt{a_n^2 - b_n^2} \ln\left(\left(a_n + \sqrt{a_n^2 - b_n^2}\right)/b_n\right),$$

e si ricorda che

$$\arccos h(x) \doteq \ln\left(x + \sqrt{x^2 - 1}\right),$$

si che

$$\arccos h(an/bn) = \ln\left(\left(a_n + \sqrt{a_n^2 - b_n^2}\right)/b_n\right)$$

e, infine,

$$OPH = (1/2)b_n\sqrt{a_n^2 - b_n^2} \arccos(a_n/b_n).$$

3. Interpretazione geometrica della successione ricorsiva in $\mathbb{R}^3$

Diamo, in questo paragrafo, una interpretazione geometrica della successione ricorsiva 1) , nei due casi di $a_0 < b_0$, $a_0 > b_0$, facendo riferimento allo spazio $\mathbb{R}^3$. Per ciò, consideriamo in $\mathbb{R}^3$ due superficie Σ_1 , Σ_2 di equazioni cartesiane, rispettivamente:

$$\Sigma_1 : \quad \sigma_1 = \frac{x+y}{2}$$

$$\Sigma_2 : \quad \sigma_2 = \sqrt{x \cdot y}$$

definite entrambe in $D = \{ (x; y) \in \mathbb{R}^2 / x > 0, y > 0 \}$: Sia $P_0(a_0; b_0) \in D$ il punto iniziale di una successione di punti di D ottenuti, a due a due, a partire da P_0 iterando i seguenti quattro passi: 1) si determina su Σ_1 il punto $Q_1(a_0; b_0; z_1)$, con $z_1 = \sigma_1(P_0)$; 2) si torna in D per considerare il punto $P_1(z_1; b_0)$, ottenuto da P_0 sostituendo alla sua ascissa la quota di Q_1 ; 3) ciò fatto, si considera in Σ_2 il punto $Q_2(z_1; b_0; z_2)$, essendo z_2 immagine di P_1 secondo σ_2 , cioè $z_2 = \sigma_2(z_1; b_0)$; 4) infine si individua, ancora in D , il punto $P_2(z_1; z_2)$ ottenuto da P_1 sostituendo la quota di Q_2 all'ordinata dello stesso P_1 ; questo processo di ricerca di punti, a partire da P_0 ci permette di passare da un punto P_k di D a un punto Q_{k+1} di Σ_1 se k è pari o a un punto Σ_2 se k è dispari, per poi tornare a un punto P_{k+1} di D ottenuto da P_k , sostituendo la quota di Q_{k+1} all'ascissa di P_k se $Q_{k+1} \in \Sigma_1$, all'ordinata di P_k se $Q_{k+1} \in \Sigma_2$. Questo processo iterativo porta a una successione di punti P_n di D , tendenti, al tendere di n a $+\infty$ al punto limite $P_l(l; l)$, essendo il valore di l dato dal secondo membro dell'espressione 7*) ovvero 7**), a seconda che P_0 appartenga all'ottante, sottoinsieme di D con $a_0 < b_0$, oppure $a_0 > b_0$, come appare evidente se si tiene presente il procedimento di individuazione dei punti $P_n \in D$, mediante le superficie Σ_1 , Σ_2 e le successioni 1). E' interessante notare che anche i punti Q_n su Σ_1 e Σ_2 , alternativamente, tendono al punto limite $Q_l(l; l; l)$ giacente sulla retta, luogo geometrico dei punti di tangenza della Σ_1 con la Σ_2 .

4. Secondo problema

Richiamiamo il problema dei poligoni regolari equivalenti: a partire da un poligono regolare di apotema a_0 e raggio r_0 , consideriamo la successione dei poligoni regolari in cui il poligono regolare successivo ha numero di lati doppio del precedente; usando la stessa simbologia del problema 1, si ottiene la seguente successione ricorsiva:

$$a_{2n} = \frac{\sqrt{2 \cdot a_n \cdot (a_n + b_n)}}{2}, \quad b_{2n} = \sqrt{a_n \cdot b_n}, \quad a_0 < r_0;$$

che, si dimostra, è anch'essa convergente.

Più in generale, si dimostra che è convergente la successione:

$$a_{n+1} = \frac{\sqrt{2 \cdot a_n \cdot (a_n + b_n)}}{2}, \quad b_{n+1} = \sqrt{a_n \cdot b_n}, \quad a_0, b_0 \in \mathbb{R}_0^+,$$

e che il limite l è dato dalla espressione 9), ovvero 10),

$$9) \quad l = \sqrt{\frac{a_0 \cdot \sqrt{b_0^2 - a_0^2}}{\arccos(a_0/b_0)}} \quad ;$$

$$10) \quad l = \sqrt{\frac{a_0 \cdot \sqrt{a_0^2 - b_0^2}}{\operatorname{arccos} h(a_0/b_0)}} \quad ,$$

a seconda che sia , rispettivamente , $a_0 < b_0$, ovvero $a_0 > b_0$, mantenendosi costante, al variare di n , nei due casi , la quantità k_n data dalla relazione 11) ovvero 12) rispettivamente:

$$11) \quad k_n = \frac{a_n \cdot \sqrt{a_n^2 - b_n^2}}{\arccos(a_n/b_n)} \quad ,$$

$$12) \quad k_n = \frac{a_n \cdot \sqrt{b_n^2 - a_n^2}}{\operatorname{arccos} h(a_n/b_n)} \quad .$$

BIBLIOGRAFIA

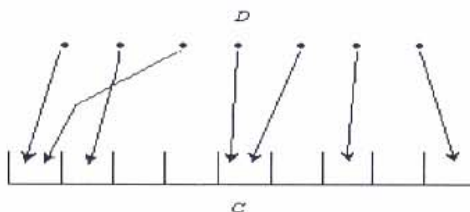
1. F. ENRIQUES, U. AMALDI, *Elementi di geometria* , vol.2, Zanichelli, Bologna, 1975.
2. E.GIUSTI, *Esercizi e complementi di analisi matematica* , vol.1, Boringhieri, Torino, 1991.

ASPETTI COMBINATORI DELLA LOGICA

Fabrizio Pascucci

I concetti dell'analisi combinatoria e soprattutto *l'angolo visuale* di questo ramo della matematica sono estremamente utili in molti campi dell'attività razionale.

Scegliamo come esempio uno dei modelli di base dell'analisi combinatoria: il modello delle biglie e delle scatole. Supponiamo di avere n biglie da mettere in m scatole. Ogni disposizione delle biglie nelle scatole corrisponde ad una funzione f di un insieme D di n elementi in un insieme C con m elementi: infatti una biglia può stare in una sola scatola e analogamente un elemento k_i di D corrisponde ad un solo elemento y_j di C , essendo $y_j = f(k_i)$.



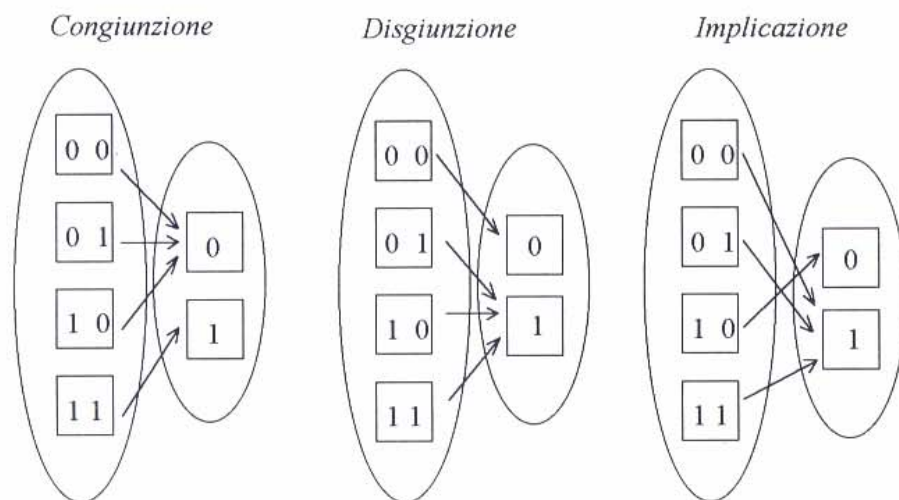
Per la prima biglia, si hanno m possibilità; passando alle successive, per ognuna ve ne sono ancora m . Il numero delle disposizioni delle biglie è quindi m^n e coincide con il numero delle funzioni dell'insieme D nell'insieme C : $|C|^{|D|}$.

Consideriamo ora due dei capitoli più noti della logica formale: il calcolo proposizionale e la teoria dei sillogismi.

(1) Gli enunciati possono essere combinati per formare enunciati complessi. Nel calcolo proposizionale si considerano solo combinazioni *vero-funzionali*, nelle quali la verità o falsità dell'enunciato composto dipende dalla verità o dalla falsità degli enunciati semplici, detti *atomici*, perché non ulteriormente scomponibili. Si possono quindi considerare *funzioni di verità* a n argomenti, $F(A_1, A_2, \dots, A_n)$, che danno il valore di verità di un enunciato composto da n enunciati atomici: ogni argomento è una variabile enunciativa, che può assumere due soli valori, identificati con 1 e 0 oppure con vero (T) e falso (F).

Il dominio D di una funzione di verità è dato da tutte le possibili assegnazioni di valori di verità alle n variabili enunciative A_i ; poiché l'insieme dei valori di A_i ha cardinalità 2, risulta $|D| = |\{0; 1\} \times \dots \times \{0; 1\}| = 2^n$. Il numero delle possibili funzioni di verità coincide con l'insieme delle funzioni di D in $\{0; 1\}$ ed è pertanto 2^{2^n} .

Nel caso che n sia 2, le funzioni sono $2^{2^2} = 16$; tra queste si trovano i connettivi più noti: la congiunzione, la disgiunzione, l'implicazione.



(2) Uno dei temi più antichi della logica è il problema dei sillogismi, che Aristotele affrontò per primo in maniera scientifica.

Ricordiamo che Aristotele riduce un enunciato qualsiasi a due componenti, i nomi e i verbi; questi esprimono la relazione che intercorre tra il nome che fa da *soggetto* (S) e quello che rappresenta il *predicato* (P). Nel sillogismo aristotelico un enunciato segue necessariamente da altri due, che rappresentano le *premesse*: le premesse contengono i *termini* della conclusione (ovvero i nomi che svolgono le funzioni di soggetto e predicato) e li collegano ad un terzo termine, detto *medio*, che costituisce il perno dell'argomentazione.

La premessa che contiene il predicato della conclusione si chiama *premessa maggiore* e l'altra, che ovviamente contiene il soggetto della conclusione, è detta *premessa minore*.

Poiché in un enunciato i nomi svolgono due funzioni, *soggetto* e *predicato*, i sillogismi si suddividono in 2×2 classi, dette *schemata* o *figure*: infatti la premessa maggiore è del tipo MP o PM , a seconda che il termine medio M vi rappresenti il soggetto o il predicato, e per ognuna delle possibilità della premessa maggiore ce ne sono due analoghe per la minore: le figure risultano allora essere

	1° fig.	2° fig.	3° fig.	4° fig.
Premessa maggiore	$M \ P$	$P \ M$	$M \ P$	$P \ M$
Premessa minore	$S \ M$	$S \ M$	$M \ S$	$M \ S$
Conclusiones	$S \ P$	$S \ P$	$S \ P$	$S \ P$

Ogni proposizione, dal punto di vista quantitativo, appartiene ad uno dei seguenti tipi:

- A, *universale affermativa* (tutti gli uomini sono mortali);
- E, *universale negativa* (nessun uomo è mortale);
- I, *particolare affermativa* (qualche uomo è mortale);
- O, *particolare negativa* (qualche uomo non è mortale).

Ad esempio, il sillogismo

Nessun eroe è codardo
 Alcuni soldati sono codardi

 Alcuni soldati non sono eroi

risulta essere della 2° figura, poiché *soldati* è il soggetto S della conclusione, *eroi* e *codardi* sono rispettivamente il predicato P della conclusione e il termine medio M ; inoltre gli enunciati che compaiono nell'argomentazione sono del tipo E, I e O.

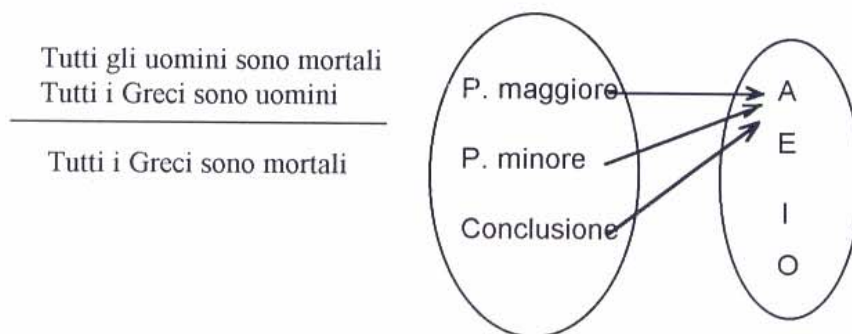
Una qualunque delle quattro figure è un insieme F di 3 elementi, gli enunciati, che poniamo in corrispondenza di un insieme di 4 valori, la classe $Q = \{A, E, I, O\}$: contando tutte le applicazioni di F in Q si danno quindi $4^3 = |Q|^{|F|}$ modi del sillogismo per ogni figura; si può anche dire che i modi sono tanti quante le disposizioni con ripetizione di classe $k=3$ di $n=4$ elementi.

I modi possibili del sillogismo sono dunque

$$|F| \times |Q|^{|F|} = 4 \times 4^3 = 4^4 = 256;$$

in realtà, i sillogismi corretti, cioè quelli che assicurano la verità della conclusione qualora siano vere le premesse, sono 6 per ciascuna figura e quindi 24 in tutto, di cui Aristotele considera solo 19. Gli schemi in figura corrispondono ai sillogismi della prima figura e della seconda, che gli Scolastici denominano, rispettivamente, *Barbara*, *Celarent* e *Baroco*:

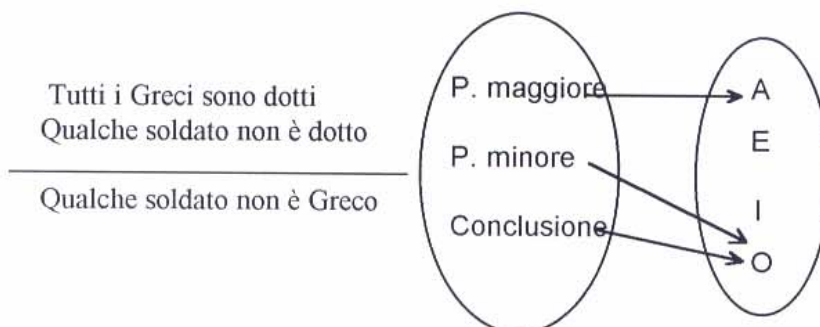
(Barbara)



(Celarent)



(Baroco)



(3) Nell'epoca della massima fioritura della civiltà araba un gruppo di astrologi arabi costruì una macchina, la *zairja*, basata su una corrispondenza tra le 28 lettere dell'alfabeto arabo e altrettanti concetti della filosofia araba; combinando le lettere secondo certi schemi si ottenevano enunciati dotati di significato, così da alimentare la leggenda della scoperta di una macchina pensante [cfr. P. Mc Corduck, *Machines Who Think*, 1979]. In uno dei suoi

tre viaggi in Africa, Raimondo Lullo (1235 - 1315) venne a conoscenza della zairja e ne fece il fulcro delle sue ricerche logiche, che culminarono nella stesura dell' *Ars Magna*. In sostanza, Lullo cercò di costruire dei "meccanismi" su base combinatoria in grado di formare enunciati accettabili e successivamente dei sillogismi validi. Elaborò quattro figure. La prima consiste in un cerchio, diviso in nove settori, detti *camere*; ciascun settore, contrassegnato da una lettera (B, C, D, E, F, G, H, J, K), corrisponde a nove nomi e ad altrettanti aggettivi (*Bonitas* e *bonus*, *Magnitudo* e *magnus*,...). Completa la figura un disco su cui è disegnata una stella che con le sue linee indica le combinazioni ammesse tra i concetti iniziali: ad esempio, *la bontà è grande* o *Dio è grande*, oppure, inversamente, *la grandezza è buona* o *la grandezza è divina*. La seconda figura, rappresentata da tre triangoli di colori diversi, serve a scegliere alcuni tra gli enunciati ottenuti mediante la prima, seguendo certi criteri (differenza, concordanza, contraddizione, ...). La terza figura, prendendo come "input" due enunciati produce il termine medio; la simulazione del procedimento sillogistico, che Lullo intende come processo di combinazione e sostituzione dei termini è finalmente portata a compimento dalla quarta figura. Questa è un sistema di tre dischi divisi in nove *camere*, di cui solo quello più esterno resta fermo interno mentre gli altri hanno un certo numero di rotazioni ammesse, corrispondenti ai sillogismi ottenibili.

Certamente i risultati di Lullo sono modesti e non sono esenti da obiezioni. Tuttavia le idee di Lullo affascinarono molti pensatori, come Leibniz e Hobbes, e, sia pur indirettamente Giuseppe Peano. Moritz Cantor ebbe a dire a proposito dell'opera di Lullo "....un miscuglio di logica, follia cabalistica e follia personale, in cui cadono, non si sa come, dei granelli di salutare buon senso."

Conclusioni. Abbiamo mostrato come uno strumento tipico dell'analisi combinatoria trovi applicazione in un campo apparentemente lontano da quello disciplinare e come più in generale sia suggestivo e fertile l'approccio combinatorio.

BIBLIOGRAFIA

1. M.CERASOLI, F. EUGENI, M. PROTASI. *Elementi di matematica discreta*, Zanichelli.
2. R. BLANCHE'. *La logica e la sua storia*. Ubaldini Editore.
3. L. GEYMONAT. *Storia del pensiero filosofico e scientifico. L'antichità e il medioevo*. Garzanti.
4. E. G. OMODEO. *L'automazione della sillogistica*. Le Scienze, Ottobre 1986, n. 218.

UN MODELLO MATEMATICO PER LA VALUTAZIONE DI FATTIBILITÀ DEI PROGRAMMI INTEGRATI D'INTERVENTO URBANO

Antonio Maturo*, Barbara Ferri*

SUNTO - Si è inteso interpretare in termini matematici un modello per la valutazione dell'effettiva realizzabilità dei Programmi di Riqualificazione e Recupero urbano (L.179/92 e L.493/93).

Tale modello è stato definito assumendo dapprima una serie di categorie di *fattibilità* che riflettono le varie componenti del benessere sociale a cui un progetto urbano deve tendere, poi individuando una serie di *criteri* che, opportunamente ponderati, riflettono le varie esigenze da soddisfare con la realizzazione di un intervento, e infine studiando gli effetti che l'opera progettata - se attuata - produrrà su tale complesso di fattori. La funzione del modello è infatti quella di verificare in che misura un progetto d'intervento urbano può garantire il perseguimento di una gamma di obiettivi di tipo economico, psicologico, fisico, culturale, urbanistico, tecnologico.

Si mostra che l'impostazione del modello in oggetto è riconducibile alla teoria delle *classificazioni fuzzy*.

ABSTRACT - In this paper we have considered a mathematical model to appraise the real feasibility of the last type of urban project. In particular we have tried to find some previsions about the effects produced by that kind of project in the urban contest.

This model is founded on:

* Dipartimento di Scienze, Storia dell'Architettura e Restauro, Viale Pindaro 42, 65127, Pescara, ITALY

- some categories of *feasibility*; these classes represent the various aspects of social welfare pursued by an urban project;
- some weighed *objects* that represent the different needs of the community whom the project is intended to;
- a verification about the attainment of aforesaid objects by the tested project.

We have noticed that the statement of this model has common something with the theory of *fuzzy sets*.

1- Posizione del problema

1.1. Il problema urbanistico

La qualità e le dimensioni dei problemi che caratterizzano le grandi conurbazioni testimoniano che la città moderna ha raggiunto una complessità difficilmente governabile.

Il genere e la qualità di infrastrutture necessarie alla vita delle società contemporanee sono notevolmente aumentate rispetto ai decenni passati, così come la congestione funzionale degli ambiti urbani, il disagio ambientale, l'espansione del recupero spontaneo; emerge una nuova domanda di centro storico e affiora l'esigenza del riuso di aree urbane aventi grandi potenzialità di incidenza urbanistica, vedendo nel riuso una risposta alla domanda di qualità della città. A questi nuovi fenomeni legati alla trasformazione del territorio si aggiunga, da una parte la necessità di rendere "trasparenti ed equi" i processi di decisione, dall'altra l'esigenza di un più mirato impiego delle risorse per il continuo ridursi di finanziamenti pubblici.

A fronte di tale complessità di fattori si ritiene ormai irrinunciabile il ricorso a nuovi strumenti di governo del territorio, in modo da travalicare la rigidità del piano regolatore generale che tende a fissare lo sviluppo della città in un'immagine statica e priva di indicazioni sulla fattibilità degli interventi in esso previsti.

Si avverte dunque l'esigenza di organiche e vincolanti procedure di valutazione che consentano di giudicare la fattibilità degli interventi sin dalle prime fasi di formazione del progetto. Queste stesse fasi sono state definite dalla legge Merloni bis del giugno '95 col nome di "*progetto preliminare*" che, secondo tale norma, deve consistere in schemi grafici volti a:

- definire il quadro delle esigenze da soddisfare;
- individuare le caratteristiche spaziali, tipologiche, funzionali e tecnologiche dei lavori da realizzare.

La legge prescrive che a questo stadio si effettuino le verifiche della *fattibilità*, l'esame dei profili di impatto ambientale, l'analisi della conformità agli strumenti urbanistici, la valutazione indicativa della spesa.

E' soprattutto nella fase iniziale che diventa strategicamente rilevante raggiungere i "punti di equilibrio" fra *obiettivi, risorse, vincoli* che possono compromettere l'esito dell'intervento.

L'attuale necessità concreta è di individuare, prevedere, verificare ogni rapporto oggetto-contesto, comprendere la totalità dei fattori implicati, in maniera da *controllare al meglio la complessità del processo progettuale*.

L'obiettivo del progettista, allora, deve essere quello di conoscere, in una specie di simulazione tipo realtà virtuale, gli effetti che l'idea progettuale, se concretizzata attraverso la realizzazione dell'opera, produrrà. In tal senso, per le opere di interesse collettivo, si ritiene che l'architetto moderno *non possa più limitarsi alla estrinsecazione grafica* delle sole sue idee, ma debba sentirsi chiamato a tradurre in opere il variegato complesso di richieste, interessi, fabbisogni e attese della collettività, nonché più in generale, il rispetto di vincoli preesistenti.

A tal fine il progettista non potrà più lavorare nel ristretto ambito del suo studio, ma dovrà essere parallelamente coadiuvato da una organizzazione che, attraverso forme di indagine scientificamente comprovate, gli permetta il controllo in tempo reale degli effetti della sua progettazione. In tal modo l'intero processo progettuale si configura come "processo decisionale in vista di alcuni obiettivi"; esso è dunque notevolmente più ampio ed impegnativo di quello connotato dalla graficità.

1.2. I programmi integrati

La difficoltà gestionale degli interventi di recupero ha spinto i diversi soggetti istituzionali alla ricerca di nuove soluzioni che affrontassero la complessità del problema inerente la riqualificazione urbana e la valorizzazione ambientale e culturale degli insediamenti.

La più recente tipologia di "progetto urbano" è rappresentata dai Programmi di Riqualificazione Urbana introdotti dalla L.179/92 e dai Programmi di Recupero Urbano della L.493/93. Tali strumenti sono l'espressione di un nuovo modo di intendere i progetti urbani poiché articolano le previsioni di piano in azioni urbanistiche di differente complessità, mettendo in gioco il ruolo di tutti gli "attori urbani".

I caratteri dei programmi suddetti si avvicinano molto a quelli di una progettazione "esecutiva" di un ambito urbano.

Il legislatore ha inteso, con la 179/92, affrontare un insieme di questioni con lo scopo esplicito di migliorare vari aspetti della gestione urbana: accanto ai temi della programmazione, dei contributi, delle locazioni e del recupero ha preso forma il tema centrale della riqualificazione urbana.

Tra i caratteri dei Programmi Integrati si trova, infatti, quello della necessaria *dimensione "strategica"*, ossia della grande incidenza sulla riorganizzazione urbana, al fine di accrescere le dimensioni funzionali e relazionali dei luoghi urbani attraverso interventi di "ricucitura"; l'interesse non è più solo per gli alloggi, ma per i servizi, le infrastrutture, le attività produttive e commerciali.

Le caratteristiche che definiscono un Programma Integrato sono:

- a) intersettorialità di funzioni;
- b) pluralità delle tipologie d'intervento: dalla manutenzione alla nuova edificazione, dal restauro alla ristrutturazione, dagli interventi edilizi a quelli urbanizzativi;
- c) pluralità di operatori;
- d) qualità urbana auspicata;
- e) modifica del ruolo del PRG, posto come norma-quadro, quindi come strumento di direttive e di orientamento più che di vincolo;
- f) modifica del ruolo dell'ente pubblico che, da motore e decisore unico, passa ad un ruolo più democratico di garante del pubblico interesse, nonché di controllo;
- g) costituzione di una forma societaria mista fra soggetti pubblici e privati, al fine di attrarre una pluralità di risorse;
- h) presenza di uno *studio di fattibilità* tecnico-economica degli interventi, al fine di rendere più efficaci i finanziamenti pubblici;
- i) "complessità", per le difficoltà imposte dai fatti urbani. Anche il quadro procedurale dei P.I. è complesso : si tratta di rendere "trasparente" la modalità di erogazione e le destinazioni di un consistente flusso di denaro pubblico, ma si tratta anche di investire sia gli aspetti di *gestione del territorio* (laddove si incide sulla strumentazione urbanistica), sia quelli di *programmazione* e di *spesa* , sia quelli di tipo *contrattuale* (si prevedono nuove convenzioni tutte da sperimentare: accordi di programma, atti d'obbligo, mandati, conferenze di servizi,...)

1.3. Caratteristiche qualitative del modello

Si è voluto formulare uno schema oggettivamente valido per affrontare la generalità dei casi relativi al tema del recupero e della riqualificazione urbana.

La funzione del modello è quella di dirigere le osservazioni sul contesto del progetto, controllare la validità di certe scelte, formulare ipotesi di intervento "ottimale" o comunque soddisfacente rispetto a determinati parametri di giudizio. Il modello proposto è organizzato in modo da poter essere strumento di valutazione "in itinere" del progetto, per permettere ai progettisti di perseguire gli obiettivi più importanti durante la fase progettuale stessa.

Il modello è basato sui concetti di *fattibilità del progetto* e di *obiettivi, o criteri*, che si intendono perseguire con la realizzazione dell'opera.

Si è ritenuto opportuno fare un po' di chiarezza su cosa significhi "studiare la fattibilità" di un progetto, poichè a tale espressione non corrisponde tuttora una definizione universalmente accettata.

L'intento di base dello studio di fattibilità del progetto deve essere quello di verificare anticipatamente le convergenze delle forze in campo, tenendo conto delle reali condizioni oggettive del contesto d'intervento. Lo studio di fattibilità può essere identificato come "guida", come "schema procedurale" o "griglia metodologica" per identificare, comprendere e descrivere i vari esiti ed effetti che potrebbero derivare dalla realizzazione di un progetto, analizzando tutte le dimensioni suscettibili d'impatto. Così, ad una fase di "analisi propedeutica" al progetto, dovrebbe far seguito una fase più mirata, tesa ad elevare la concretezza dello studio e a meglio valutare l'intervento in termini di qualità urbana e resa socio-economica, attraverso la costituzione di un *modello*.

Tale modello si fonda sul principio che lo studio di fattibilità dei programmi di recupero e riqualificazione urbana si configurino come scelte "complesse" in cui è necessario perseguire contemporaneamente più obiettivi eterogenei e talvolta conflittuali. Si è ritenuto perciò di ricorrere alle tecniche di valutazione *multicriterio* che consentono di rendere lo studio "onnicomprensivo", considerando pure le dimensioni "extraeconomiche" di un intervento, nonché le ampie ripercussioni che possono aversi anche sugli individui non direttamente interessati dall'opera progettata. Le *esternalità* e gli *intangibili* sono infatti particolarmente rilevanti negli interventi di recupero e riqualificazione urbana.

Per esprimere tale circostanza è parso necessario assumere nel modello una certa serie di *fattibilità* che riflettono ed esprimono tutte le componenti in gioco, e una certa serie di *obiettivi* - o *criteri* - che, in seno a ciascuna categoria di fattibilità, riflettono ed esprimono i bisogni o le richieste della collettività.

Sia le fattibilità che i criteri possono essere *ponderati* opportunamente per tener conto della importanza riconosciuta alle diverse esigenze. In tal caso il peso di ciascuna fattibilità ne caratterizza l'importanza relativa nell'ambito di una prefigurata tipologia progettuale d'intervento in cui siano state prefissate determinate priorità. I pesi dei criteri rappresentano invece le misure del grado di importanza dei vari obiettivi. L'attribuzione dei pesi deve avvenire sulla base del parere di tutti i soggetti del piano: promotori, attivi, o semplicemente coinvolti.

La valutazione diventa allora *l'operazione con cui si riesce ad esprimere il grado fino al quale i diversi obiettivi sono soddisfatti dalle caratteristiche del progetto*.

In un programma di riqualificazione urbana gli obiettivi si riferiscono oltre che alla eliminazione del degrado sociale ed economico di un'area, anche al soddisfacimento di esigenze psicologiche, di tutela dei valori storico-artistici, di tutela delle qualità estetico-visive. La formazione di un programma d'intervento urbano deve configurarsi come processo di sintesi il più possibile "globale" delle diverse componenti in gioco.

A tal fine sono state considerate sette categorie di fattibilità:

1. F. AMBIENTALE, intesa come valutazione degli effetti del progetto sull'ambiente naturale (fisico e paesaggistico) e sull'ambiente costruito. A tale riguardo bisogna considerare le caratteristiche morfologiche del sito d'intervento, le caratteristiche del terreno in senso fisico-meccanico, geologico, idrogeologico, pedologico, gli elementi caratterizzanti il sistema urbano, gli elementi fisici condizionanti gli interventi sulla zona.
2. F. ESTETICO-CULTURALE, intesa come valutazione degli effetti del progetto sulla tutela degli interessi storici, artistici, archeologici. Aspetto sostanziale di tale fattibilità è l'analisi di "come" l'intervento si rapporta alle preesistenze artificiali, alla loro importanza storica e architettonica, al fine di garantire il rispetto dei valori (urbani e artistici) suddetti. Elementi da sottoporre allo studio sono: aspetti culturali locali, valori simbolici del sito. La lettura di "costruzione storica" dell'ambiente è condotta per far sì che l'intervento si concili con la cultura e la storia che hanno dato forma e senso alla città.
3. F. ECONOMICA, intesa come valutazione tesa ad accertare che l'intervento contribuisca al miglioramento della base economica urbana. Utilizzando il miglioramento dei valori funzionali dei manufatti e della loro qualità fisica è possibile determinare impatti positivi sulle attività, stimolando la localizzazione di nuove iniziative. Questo significa che l'attività di conservazione non deve essere solo finalizzata al miglioramento dello scenario fisico.
4. F. FINANZIARIA, intesa come verifica della provvista finanziaria disponibile per la realizzazione del programma, attraverso l'individuazione di tutti i canali finanziari attivabili.
5. F. TECNICA, intesa come studio degli effetti del progetto su alcune componenti tecniche: vincoli tecnico-costruttivi, vincoli derivanti dalla strumentazione urbanistica, problematiche infrastrutturali, rapporti funzionali e strutturali tra l'opera e il contesto urbano, difficoltà eventuali relative alla organizzazione del cantiere.
6. F. SOCIALE, riguarda il perseguimento di obiettivi che investono i soggetti direttamente interessati dal progetto. A tal fine bisogna considerare gli eventuali effetti di risanamento urbano complessivo: il miglioramento della qualità insediativa, l'eliminazione del degrado sociale, l'incremento del tasso occupazionale, la tutela delle categorie deboli, eventuali domande non soddisfatte

7. F. PROCEDURALE, volta a considerare sostanzialmente tre aspetti dell'intervento: l'*aspetto regolamentare*, ossia la congruenza del programma rispetto alle prescrizioni vigenti; l'*aspetto contrattuale*, relativo alla compresenza di differenti attori legati da "contratti"; l'*aspetto operativo*, relativo alle modalità di realizzazione dell'opera (disponibilità degli immobili, affidamento dei lavori, aspetti gestionali).

Per ognuna delle sette fattibilità, si valuta l' "efficacia" del progetto relativamente ad ogni criterio in essa compreso; si individua cioè il grado con cui il progetto soddisfa ciascun obiettivo rispetto alla fattibilità considerata.

Si è ritenuto inoltre assegnare ad alcuni criteri l'etichetta di "essenzialità" e ad altri di "secondarietà". La differenziazione suddetta, effettuata in base alla importanza prioritaria che ad un certo criterio viene assegnata dal complesso di normative sui P.I., è dettata dalla esigenza di imporre che il progetto persegua i criteri essenziali in misura almeno sufficiente ed è stata ottenuta assegnando opportunamente dei pesi.

Si sono poi introdotti nell'insieme dei progetti una relazione di *Preordine* ed una funzione detta *Valutazione Numerica* o *Utilità Globale* dei progetti.

Attraverso l'assunzione di una *matrice di soglia* e di due valori detti rispettivamente *valore scalare di soglia* e *valore di apprezzamento*, da confrontare, rispettivamente, con la matrice che individua i gradi con cui il progetto soddisfa gli obiettivi rispetto alle fattibilità considerate e con il valore dell'Utilità Globale del progetto esaminato, sarà possibile stabilire in che misura il progetto risulta fattibile.

Viene infine introdotto il concetto di progetto *quasi fattibile ma molto valido*.

2 - Un modello matematico di "ricoprimento fuzzy" per la valutazione di un progetto

2.1. Ricoprimenti fuzzy di un insieme Ω e relazioni di preordine fra progetti

Una *proprietà P in senso usuale* o *proprietà hard* o *proprietà crisp* su un insieme Ω si può descrivere come un' affermazione o *vera* o *falsa* per ogni oggetto $x \in \Omega$. Poniamo $P(x) = 1$ se per x la proprietà è vera e $P(x) = 0$ se per x la proprietà è falsa.

In altre parole P può essere considerata come una funzione definita in Ω ed a valori in $\{0,1\}$. I sottoinsiemi di Ω

$$P_0 = \{x \in \Omega : P(x) = 0\}, \quad P_1 = \{x \in \Omega : P(x) = 1\}$$

si chiamano, rispettivamente, *parte falsa* e *parte vera* di Ω determinate dalla proprietà P .

Una proprietà P *sfumata* o *sfocata* o *fuzzy* su un insieme Ω si può descrivere come un' affermazione o *vera* o *falsa* o *parzialmente vera* per ogni oggetto $x \in \Omega$. Poniamo $P(x) = 1$ se per x la proprietà è vera, $P(x) = 0$ se per x la proprietà è falsa e $P(x) = a$, con $0 < a < 1$, se la proprietà è parzialmente vera e abbiamo una misura del *grado di verità* della proprietà data dal numero a .

In altre parole P può essere considerata come una funzione definita in Ω ed a valori in $[0,1]$. Per ogni $a \in [0,1]$, essa determina il sottoinsieme di Ω

$$P_a = \{x \in \Omega : P(x) = a\},$$

detto *parte di livello* a determinata da P .

L'unione dei sottoinsiemi P_a , con $a \in (0,1)$ si chiama *parte sfumata* o *sfocata* o *fuzzy* determinata dalla proprietà P . Essa è vuota se e solo se P è una proprietà *hard*.

In tale ordine di idee diamo le seguenti definizioni:

Definizione 1. Diciamo *fuzzy set* o *insieme sfocato* con insieme universo Ω ogni funzione

$$f : \Omega \rightarrow [0,1]. \quad (1)$$

Per ogni $x \in \Omega$, $f(x)$ si dice *grado di appartenenza* di x ad f .

I sottoinsiemi di Ω , $V(f) = f^{-1}(1)$, $F(f) = f^{-1}(0)$, $I(f) = f^{-1}((0,1))$, si dicono, rispettivamente, *parte vera*, *parte falsa* e *parte sfumata* di f .

Sia D la famiglia dei sottoinsiemi finiti o numerabili di Ω . Diciamo *cardinalità* di f il numero

$$\text{card}(f) = \sup \left\{ \sum_{x \in I} f(x), I \in D \right\}. \quad (2)$$

Un fuzzy set si dice *hard set* o *insieme usuale* o *crisp set* se risulta $f(\Omega) \subseteq \{0,1\}$, ossia se $I(f) = \emptyset$.

Osservazione 2. Un fuzzy set è individuato dalla coppia formata da $V(f)$ e dalla famiglia $\{f^{-1}(a)\}_{a \in (0,1)}$. Se il fuzzy set è un crisp set allora esso è individuato semplicemente dall'insieme $V(f)$ ed in pratica viene identificato con esso.

Definizione 3. Sia $F = \{f_h\}_{h \in H}$ una famiglia di fuzzy set con insieme universo Ω e insieme di indici H . Sia T la famiglia dei sottoinsiemi finiti o numerabili di H . Diciamo *grado di appartenenza* di x ad F il numero

$$\text{gr}(x) = \sup \left\{ \sum_{h \in J} f_h(x), J \in T \right\}. \quad (3)$$

Corollario 4. Supponiamo che sia H che Ω siano finiti o numerabili. Allora risulta:

$$\forall x \in \Omega, \quad \text{gr}(x) = \sum_{h \in H} f_h(x); \quad (4)$$

$$\forall h \in H, \quad \text{card}(f_h) = \sum_{x \in \Omega} f_h(x); \quad (5)$$

$$\sum_{x \in \Omega} \text{gr}(x) = \sum_{h \in H} \text{card}(f_h). \quad (6)$$

Definizione 5. Sia H un insieme e sia $F = \{f_h\}_{h \in H}$ una famiglia di fuzzy set con insieme universo Ω e insieme di indici H .

Gli insiemi $\text{DR}(F) = \{x \in \Omega : \text{gr}(x) > 0\}$ e $\text{FR}(F) = \{x \in \Omega : \text{gr}(x) \geq 1\}$ si dicono, rispettivamente, *parte debolmente ricoperta* di Ω e *parte fortemente ricoperta* di Ω . La F si dice *fuzzy ricoprimento parziale* di Ω se $\text{DR}(F) \subset \Omega$, si dice *fuzzy ricoprimento debole* se $\text{DR}(F) = \Omega$, si dice *fuzzy ricoprimento forte* se $\text{FR}(F) = \Omega$.

Definizione 6. Sia H un insieme e sia $F = \{f_h\}_{h \in H}$ una famiglia di fuzzy set con insieme universo Ω e insieme di indici H .

Diciamo che F è una *fuzzy partizione parziale, debole o forte* di Ω se è, rispettivamente, un *fuzzy ricoprimento parziale, debole o forte* di Ω ed inoltre valgono le seguenti proprietà

$$(P1) \quad \forall h \in H, \exists x \in \Omega : f_h(x) > 0;$$

$$(P2D) \quad \forall x \in \Omega, \text{gr}(x) \leq 1.$$

Corollario 7. La F è una *fuzzy partizione forte* di Ω se e solo se valgono la (P1) e la

$$(P2F) \quad \forall x \in \Omega, \text{gr}(x) = 1.$$

Corollario 8. Siano H ed Ω due insiemi finiti o numerabili e sia $F = \{f_h\}_{h \in H}$ una famiglia di fuzzy set con insieme universo Ω e insieme di indici H . Allora F è una fuzzy partizione forte di Ω se e solo se

$$(P1A) \quad \forall h \in H, \quad \sum_{x \in \Omega} f_h(x) > 0;$$

$$(P2A) \quad \forall x \in \Omega, \quad \sum_{h \in H} f_h(x) = 1$$

Definizione 9. Sia f un fuzzy set con insieme universo Ω . Per ogni fuzzy set g con insieme universo Ω il prodotto

$$f \bullet g : x \in \Omega \rightarrow f(x) g(x) \quad (7)$$

si dice *parte* di f *trasmessa* da g ed il prodotto

$$f \bullet (1-g) : x \in \Omega \rightarrow f(x) [1 - g(x)] \quad (8)$$

si dice *parte* di f *assorbita* da g . La coppia $(g, \bullet)$ si dice *filtro* rispetto ad f .

Definizione 10. Sia $F = \{f_h\}_{h \in H}$ una famiglia di fuzzy set con insieme universo Ω e insieme di indici H . Una famiglia $G = \{g_h\}_{h \in H}$ di fuzzy set con insieme universo Ω e insieme di indici H si dice *filtro* rispetto a F se, $\forall h \in H$, g_h è un filtro rispetto a f_h .

Si dice *parte* di F *trasmessa* da G la famiglia

$$T(F, G) = \{f_h \bullet g_h\}_{h \in H} \quad (9)$$

e si dice *parte* di F *assorbita* da G la famiglia

$$A(F, G) = \{f_h \bullet (1 - g_h)\}_{h \in H}. \quad (10)$$

Definizione 11. Sia $F^I = \{F^\alpha\}_{\alpha \in I}$ un insieme, con insieme di indici I , di famiglie di fuzzy set $F^\alpha = \{f_h^\alpha\}_{h \in H}$ con insieme universo Ω e insieme di indici H e sia $S = \{s_h\}_{h \in H}$ una famiglia di fuzzy set con insieme universo Ω e insieme di indici H .

Chiamiamo *operatore di soglia* definito in F^I con *matrice di soglia* S , la terna (F^I, S, φ) dove φ è la famiglia di funzioni $\{\varphi^\alpha\}_{\alpha \in I}$ con φ^α definita in $\Omega \times H$ e tale che, $\forall (x, h) \in \Omega \times H$, $\varphi^\alpha(x, h) = 1$ se $f_h^\alpha(x) \geq s_h(x)$ e $\varphi^\alpha(x, h) = 0$ se $f_h^\alpha(x) < s_h(x)$.

Supponiamo da ora in poi che sia H che Ω siano finiti. Indichiamo con m il numero di elementi di Ω e con n il numero di elementi di H .

Definizione 12. Sia $F = \{f_h\}_{h \in H}$ una famiglia di fuzzy set con insieme universo Ω e insieme di indici H . Si dice *sistema di pesi relativi* su Ω ogni fuzzy set π con insieme universo Ω e cardinalità uguale ad 1. Si dice *sistema di pesi relativi* su F ogni fuzzy set λ con insieme universo H e cardinalità uguale ad 1.

Definizione 13. Sia $F^I = \{F^\alpha\}_{\alpha \in I}$ un insieme, con insieme di indici I , di famiglie di fuzzy set $F^\alpha = \{f_h^\alpha\}_{h \in H}$ con insieme universo Ω e insieme di indici H e siano $S = \{s_h\}_{h \in H}$, $T = \{t_h\}_{h \in H}$ due famiglie di fuzzy set tali che, $\forall x \in \Omega, \forall h \in H, \forall \alpha \in I$,

$$\max \{s_h(x), f_h^\alpha(x)\} \leq t_h(x). \quad (11)$$

Sia inoltre $\lambda = \{\lambda_h\}_{h \in H}$ un sistema di pesi relativi su H e $\pi = \{\pi_x\}_{x \in \Omega}$ un sistema di pesi relativi su Ω .

Chiamiamo *struttura di progetti* con tetto T e soglia S la sestupla $(F^I, S, \varphi, T, \lambda, \pi)$, con (F^I, S, φ) operatore di soglia definito in F^I e matrice di soglia S . Chiamiamo *progetti* gli elementi di F^I .

Chiamiamo inoltre:

(a) *media* di F^α il numero $m(F^\alpha) = \sum_{h \in H} \sum_{x \in \Omega} \pi_x \lambda_h f_h^\alpha(x), \forall \alpha \in I$;

(b) *media di soglia* il numero $m(S) = \sum_{h \in H} \sum_{x \in \Omega} \pi_x \lambda_h s_h(x)$;

(c) *media di tetto* il numero $m(T) = \sum_{h \in H} \sum_{x \in \Omega} \pi_x \lambda_h t_h(x)$;

(d) *carenza media* di F^α rispetto alla soglia S il numero

$$c_s(F^\alpha) = \sum_{f_h^\alpha(x) \leq s_h(x)} \pi_x \lambda_h (s_h(x) - f_h^\alpha(x)); \quad (12)$$

(e) *eccedenza media* di F^α rispetto alla soglia S il numero

$$e_s(F^\alpha) = \sum_{f_h^\alpha(x) \geq s_h(x)} \pi_x \lambda_h (f_h^\alpha(x) - s_h(x)); \quad (13)$$

(f) *soddisfacimento scalare di F^α rispetto a S* il numero

$$d_s(F^\alpha) = m(F^\alpha) - m(S).$$

Osservazione 14. Ogni progetto F^α si può rappresentare come punto $P(F^\alpha) = (e_s(F^\alpha), c_s(F^\alpha))$, del piano (e, c) , dove e indica l'eccedenza e c la carenza di F^α . Precisamente, essendo

$$0 \leq e_s(F^\alpha) \leq m(T) - m(S), \quad 0 \leq c_s(F^\alpha) \leq m(S), \quad (14)$$

$P(F^\alpha)$ appartiene al rettangolo $R = [0, m(T) - m(S)] \times [0, m(S)]$.

Definizione 15. Diciamo che due progetti sono *(e,c)-equivalenti* se sono rappresentati da uno stesso punto.

Diciamo che due progetti sono *scalarmente equivalenti* se hanno lo stesso soddisfacimento scalare.

Notiamo che risulta $e_s(F^\alpha) - c_s(F^\alpha) = d_s(F^\alpha)$. Segue la

Proposizione 16. Due progetti sono scalarmente equivalenti se e solo se sono rappresentati da punti che si trovano su una stessa retta parallela alla bisettrice degli assi.

Definizione 17. Dati due progetti F^α e F^β , $\alpha \in I$, $\beta \in I$, diciamo che F^α è *(e,c)-suvvalente* rispetto a F^β ovvero che F^β è *(e,c)-prevalente* rispetto a F^α e scriviamo $F^\alpha \leq F^\beta$ se risulta $e_s(F^\alpha) \leq e_s(F^\beta)$ e $c_s(F^\alpha) \geq c_s(F^\beta)$.

Supponiamo, da ora in poi, che F^I sia finito e che ad esso appartengano un progetto F^m , detto *massimale* tale che $F^\alpha \leq F^m$, $\forall \alpha \in I$, ed un progetto F^0 , detto *minimale* tale che $F^0 \leq F^\alpha$, $\forall \alpha \in I$.

Si può convenire, ad esempio, di assumere come massimale un *progetto ideale* F^m tale che $e_s(F^m) = m(T) - m(S)$, $c_s(F^m) = 0$.

Si può pensare inoltre di assumere come minimale un *progetto nullo* F^0 con $e_s(F^0) = 0$, $c_s(F^0) = m(S)$.

Si consideri la funzione

$$P : F^\alpha \in F^I \rightarrow P(F^\alpha) \in R. \quad (15)$$

Poniamo, $\forall P_1(e_1, c_1), P_2(e_2, c_2) \in R$,

$$P_1 \leq P_2 \Leftrightarrow (e_1 \leq e_2 \text{ e } c_1 \geq c_2). \quad (16)$$

La $\leq$ è una relazione d'ordine in R e risulta

$$F^\alpha \leq F^\beta \Leftrightarrow P(F^\alpha) \leq P(F^\beta). \quad (17)$$

Indichiamo con R^* l'immagine di F^I tramite la P e consideriamo la restrizione di $\leq$ ad R^* .

Per la (17) risulta $P(F^0) \leq P(F^\alpha) \leq P(F^m)$ per ogni progetto $F^\alpha \in F^I$. Allora, dati due progetti F^α e F^β , gli insiemi

$$M_1 = \{F^i \in F^I : P(F^i) \leq P(F^\alpha) \text{ e } P(F^i) \leq P(F^\beta)\},$$

$$M_2 = \{F^i \in F^I : P(F^\alpha) \leq P(F^i) \text{ e } P(F^\beta) \leq P(F^i)\},$$

sono non vuoti, poichè $F^0 \in M_1$ e $F^m \in M_2$.

Possiamo allora dare la seguente

Definizione 18. Per ogni coppia F^α e F^β di progetti appartenenti ad F^I , siano $P(F^\alpha) \wedge P(F^\beta)$ l'insieme dei minoranti massimali di $P(F^\alpha)$ e $P(F^\beta)$, $P(F^\alpha) \vee P(F^\beta)$ l'insieme dei maggioranti minimali di $P(F^\alpha)$ e $P(F^\beta)$.

Diciamo *iperprodotto* $\wedge$ in R^* la legge che ad ogni coppia $(P(F^\alpha), P(F^\beta))$ di elementi di R^* associa $P(F^\alpha) \wedge P(F^\beta)$ e diciamo *iperprodotto* $\vee$ in R^* la legge che ad ogni $(P(F^\alpha), P(F^\beta)) \in R^* \times R^*$ associa $P(F^\alpha) \vee P(F^\beta)$.

Definizione 19. Si dice *iperprodotto* $\wedge$ in F^I la legge che ad ogni coppia $(F^\alpha, F^\beta) \in F^I \times F^I$ associa $P^{-1}(P(F^\alpha) \wedge P(F^\beta))$ e si dice *iperprodotto* $\vee$ in F^I la legge che ad ogni $(F^\alpha, F^\beta) \in F^I \times F^I$ associa $P^{-1}(P(F^\alpha) \vee P(F^\beta))$.

Teorema 20. $(F^I, \wedge)$ soddisfa alle seguenti proprietà

$$(\wedge 1) \quad \forall F \in F^\alpha \wedge F^\beta, F \leq F^\alpha \text{ e } F \leq F^\beta;$$

$$(\wedge 2) \quad \forall F \in F^I, (F \leq F^\alpha, F \leq F^\beta) \Rightarrow \exists G \in F^\alpha \wedge F^\beta: F \leq G.$$

Teorema 21. $(F^I, \vee)$ soddisfa alle seguenti proprietà

$$(\vee 1) \quad \forall F \in F^\alpha \vee F^\beta, F^\alpha \leq F \text{ e } F^\beta \leq F;$$

$$(\vee 2) \quad \forall F \in F^I, (F^\alpha \leq F, F^\beta \leq F) \Rightarrow \exists G \in F^\alpha \vee F^\beta: G \leq F.$$

2.2. Il nostro modello come partizione fuzzy con "oggetti pesati"

Indichiamo con $\Omega = \{x_1, x_2, \dots, x_m\}$ l'insieme delle fattibilità e con $C = \{c_1, c_2, \dots, c_n\}$ quello dei criteri, o obiettivi, che possono essere perseguiti dalle fattibilità stesse.

Per ogni $x_i \in \Omega$ e per ogni $c_j \in C$ indichiamo con p_{ij} un numero appartenente all'intervallo $[0,1]$ che rappresenta il *grado di importanza relativa* del criterio c_j rispetto alla fattibilità x_i , detto anche *grado di appartenenza* di x_i a c_j o anche *misura* di c_j rispetto ad x_i .

Ogni criterio c_i può essere considerato come un fuzzy set con insieme universo Ω e viene rappresentato dal vettore colonna

$$c_j = \begin{bmatrix} p_{1j} \\ p_{2j} \\ \vdots \\ p_{mj} \end{bmatrix}$$

Sia $H = \{ 1, 2, \dots, n \}$. La famiglia dei fuzzy set $\{c_j\}_{j \in H}$ con insieme universo Ω è allora rappresentato dalla seguente matrice detta *matrice caratteristica del modello matematico* o *matrice di fattibilità*.

$$M = \begin{array}{c} \begin{array}{ccccccc} & c_1 & c_2 & \dots & c_j & \dots & c_n \\ \begin{array}{c} x_1 \\ x_2 \\ \vdots \\ x_i \\ \vdots \\ x_m \end{array} & \begin{array}{c} p_{11} \\ p_{12} \\ \vdots \\ p_{ij} \\ \vdots \\ p_{1n} \end{array} \end{array} \end{array}$$

Dal corollario 8 di 2.1. seguono le relazioni

$$\forall x_i \in \Omega, \text{ gr } (x_i) = \sum_{j=1}^n p_{ij}; \quad (18)$$

$$\forall c_j \in C, \text{ card } (c_j) = \sum_{i=1}^m p_{ij}. \quad (19)$$

Chiamiamo *cardinalità* di C il numero

$$\text{card } (C) = \sum_{j=1}^n \text{card } (c_j). \quad (20)$$

Dalla (6) si ricava che

$$\text{card}(C) = \sum_{i=1}^m \text{gr}(x_i). \quad (21)$$

La condizione $\sum_{j=1}^n p_{ij} = 1$ sulle righe della matrice di fattibilità

equivale alla (PF2) per cui, dato che ogni criterio ha cardinalità positiva, la famiglia C risulta essere una partizione forte di Ω .

Dato un progetto F , per ogni $x_i \in \Omega$ e per ogni $c_j \in C$ indichiamo con e_{ij} un numero appartenente all'intervallo $[0, 1]$ che rappresenta il *grado di efficacia* con cui il progetto F raggiunge l'obiettivo c_j relativamente alla fattibilità x_i .

Per ogni fissato criterio c_j l'insieme dei gradi di efficacia può essere considerato come un fuzzy set con insieme universo Ω e viene rappresentato dal vettore colonna

$$e_j = \begin{bmatrix} e_{1j} \\ e_{2j} \\ \vdots \\ e_{mj} \end{bmatrix}$$

La famiglia dei fuzzy set $\{e_j\}_{j \in H}$ con insieme universo Ω è allora rappresentato dalla seguente matrice

$$E = \begin{array}{c} \begin{array}{ccccccc} & c_1 & c_2 & \dots & c_j & \dots & c_n \\ \begin{array}{c} x_1 \\ x_2 \\ \vdots \\ x_i \\ \vdots \\ x_m \end{array} & \begin{array}{c} e_{11} \\ e_{12} \\ \vdots \\ e_{ij} \\ \vdots \\ e_{1n} \end{array} & \begin{array}{c} e_{21} \\ e_{22} \\ \vdots \\ e_{2j} \\ \vdots \\ e_{2n} \end{array} & \dots & \begin{array}{c} e_{j1} \\ e_{j2} \\ \vdots \\ e_{jj} \\ \vdots \\ e_{jn} \end{array} & \dots & \begin{array}{c} e_{n1} \\ e_{n2} \\ \vdots \\ e_{nj} \\ \vdots \\ e_{nn} \end{array} \end{array} \end{array}$$

detta *matrice di efficacia del progetto*. Essa dipende dalle caratteristiche dello specifico progetto esaminato.

La matrice E , considerata come insieme dei suoi vettori colonna $\{e_j\}_{j \in H}$ si può considerare come un filtro rispetto alla matrice M , anch'essa presa come insieme dei suoi vettori colonna $\{c_j\}_{j \in H}$.

La parte di M trasmessa da E , $T(M, E) = \{c_h \cdot e_h\}_{h \in H}$ rappresenta i *gradi di appartenenza* dei criteri c_j rispetto alle fattibilità x_i filtrati dal progetto F .

La parte di M assorbita da E , $A(M, E) = \{c_h \cdot (1 - e_h)\}_{h \in H}$ rappresenta le *perdite* dei gradi di appartenenza dei criteri c_j rispetto alle fattibilità x_i con il progetto F .

In seguito rappresentiamo il progetto F con la sua matrice $T(M, E)$, vista come insieme dei suoi vettori colonna.

Assegniamo, in base a criteri determinati da esperti, un sistema di pesi $\pi = \{\pi_j\}_{j \in H}$ ai criteri, un sistema di pesi $\lambda = \{\lambda_x\}_{x \in \Omega}$ alle fattibilità ed una matrice di soglia $S = \{s_j\}_{j \in H}$, con

$$s_j = \begin{bmatrix} s_{1j} \\ s_{2j} \\ \vdots \\ s_{mj} \end{bmatrix}$$

vettore di soglia associato al criterio c_j , tale che, $\forall i \in \{1, \dots, m\}, \forall j \in \{1, \dots, n\}$, risulti $s_{ij} \leq p_{ij}$.

Ogni s_{ij} rappresenta il minimo grado accettabile in cui il criterio c_j deve essere soddisfatto rispetto alle fattibilità x_i .

Se F^I è l'insieme dei progetti considerati viene allora ad essere definito un operatore di soglia (F^I, S, φ) ed una struttura di progetti $(F^I, S, \varphi, T, \lambda, \pi)$, con T uguale alla matrice M di fattibilità considerata come insieme dei suoi vettori colonna.

Si possono allora prendere in considerazione tutti i concetti introdotti nella definizione 13 e nelle definizioni e teoremi successivi e rappresentare i progetti come punti del piano (e, c) .

In particolare, due progetti F^α e F^β risultano essere (e, c) -equivalenti se $P(F^\alpha) = P(F^\beta)$ e la (17) fornisce una relazione di preordine nell'insieme F^I dei progetti.

2.3. Valutazioni numeriche e classificazione dei progetti

Consideriamo una struttura di progetti $(F^I, S, \varphi, T, \lambda, \pi)$ utilizzando le notazioni del paragrafo precedente, in particolare della definizione 13 e conseguenze.

Una prima valutazione numerica di un progetto F^α è data dal suo soddisfacimento scalare $d_s(F^\alpha)$. Poichè risulta

$$F^\alpha \leq F^\beta \Rightarrow d_s(F^\alpha) \leq d_s(F^\beta) \quad (22)$$

tale valutazione numerica è compatibile con la relazione di preordine $\leq$.

Spesso si assume la condizione

$$c_s(F^\alpha) = 0, \quad (23)$$

ossia che per ogni coppia (x_i, c_j) risulti $f_{ij}^\alpha \geq s_{ij}$.

In tal caso è $d_s(F^\alpha) = e_s(F^\alpha)$ e risulta

$$F^\alpha \leq F^\beta \Leftrightarrow d_s(F^\alpha) \leq d_s(F^\beta). \quad (24)$$

La relazione di preordine in F^I data dalla d_s coincide con la $\leq$.

Noi riteniamo, tuttavia, che si debbano prevedere condizioni più elastiche della (23), basate sull'idea che un progetto F^α possa avere un valore rilevante anche se non supera la soglia rispetto a qualche coppia fattibilità - criterio, ossia se esistono coppie (x_i, c_j) tali che $f_{ij}^\alpha < s_{ij}$, e che in tal caso sia opportuno in qualche modo recuperarlo facendo qualche modifica al progetto o vedendo se è possibile abbassare i valori di soglia con approfondite consultazioni con esperti o politici.

Riteniamo pertanto che si debba individuare un valore opportuno $h > 0$ e considerare oltre alla (23) la condizione più debole

$$c_s(F^\alpha) \leq h. \quad (25)$$

Inoltre pensiamo che si debba valutare scalarmente il progetto tramite una funzione reale v , definita in F^I tale che

$$F^\alpha \leq F^\beta \Rightarrow v(F^\alpha) \leq v(F^\beta). \quad (26)$$

Tale funzione può essere la d_s , ma esigenze di varia natura potrebbero far pervenire una funzione v diversa dalla d_s . In questo lavoro ci limitiamo al caso in cui $v = d_s$.

Inoltre bisogna fissare dei valori v_1 e v_2 di v con $v_1 < v_2$, detti "valore scalare di soglia" e "valore di apprezzamento". Nel caso in cui $v = d_s$ riteniamo che si debba assumere $v_1 = m(S)$.

I progetti F^α appartenenti a F^I vengono allora classificati in base ai valori $c_s(F^\alpha)$ e $v(F^\alpha)$ in 9 classi.

In base alla c_s i progetti si dicono

(C1) fattibili se $c_s = 0$

(C2) quasi fattibili se $0 < c_s \leq h$

(C3) non fattibili se $c_s > h$.

In base alla v i progetti si dicono

(V1) molto validi se $v \geq v_2$

(V2) sufficientemente validi se $v_1 \leq v < v_2$

(V3) non validi se $v < v_1$.

Se si assume la (23) si considerano solo le tre classi (C1) - (V1), (C1) - (V2) e (C1) - (V3). In particolare si trascura la classe (C2) - (V1) costituita da progetti molto validi e che superano di poco la matrice di soglia, progetti che riteniamo debbano essere valutati in maniera approfondita.

BIBLIOGRAFIA

1. A. FADINI, *Introduzione alla teoria degli insiemi sfocati*, Montefeltro Liguori Editore, 1979.
2. A. KAUFMANN, *Theory of fuzzy subsetc*, vol. I Academie Press, New York 1975
3. A. KAUFMANN, *Introduction à la théorie des sous-ensembles flous*, vol II *Applications à la semantique*, Masson, Paris 1975
4. A. KAUFMANN, *Introductions ecc.* Vol III - *Applications a la classification et à la reconossance des formes*, Masson, Paris 1975
5. F. GIRARD, *Risorse architettoniche e culturali: valutazioni e strategie di conservazione*, Franco Angeli, Milano 1987
6. F. GIRARD, *La conservazione nella pianificazione fisica*, Franco Angeli, Milano 1989
7. A. ZADETH, *Fuzzy sets*, Inform. Contr., 8, 1965
8. S. MICCOLI, *La valutazione di fattibilità nei programmi complessi d'intervento urbano*, Genio Rurale n°3, 1995

ALGEBRA ASTRATTA E INSEGNAMENTO

Luigi Cerritelli* , Fabio Mercanti*

SUNTO - Si mostra un esempio di *iperoperazione* su un cerchio, come applicazione didattica di *un'algebra astratta*.

ABSTRACT - An example of *hyperoperation* on a circle has been exhibited, as a didactic application of an *abstract algebra*.

INTRODUZIONE

0. In questa nota si propone l'uso di particolari strutture algebriche, le *iperstrutture*, per interpretare alcuni aspetti della geometria elementare con un linguaggio algebrico astratto. Senza addentrarsi nella problematica assai complessa dello studio delle iperstrutture, per le quali si rimanda a CORSINI [2], [3], BERARDI-EUGENI-INNAMORATI [1], ci si limita, in questa sede, solamente a richiamare il concetto di *iperoperazione*, che sta alla base di quello di iperstruttura. Si mostra poi un esempio di applicazione didattica, passibile di ulteriore miglioramento, del concetto di iperoperazione stessa.

I PROGRAMMI DI FRASCATI

1. Lo spirito informatore dei programmi studiati, alla fine degli anni sessanta, nei convegni promossi dalla *Commissione per l'Insegnamento*

* Politecnico di Milano, Via Bonardi 3.

della Matematica, i cosiddetti *programmi di Frascati*, era fortemente indirizzato verso la ricerca di una base matematica comune a tutti gli indirizzi di scuola secondaria superiore. Tale base veniva successivamente sviluppata in funzione delle specializzazioni triennali. In tal senso tra i primi libri scolastici innovatori ricordiamo *Matematica 1*, Zanichelli, Bologna 1971, di A. Rossi Dell'Acqua e F. Speranza. In esso, infatti, si individuava, nel linguaggio dell'algebra astratta, il fulcro attorno al quale costruire una unitarietà che descrivesse algebra e geometria euclidea con un continuo collegamento trasversale. Fu una operazione *fusionista* (D'ANIELLO[4]) nel senso più ardito della parola: una provocazione intelligente e anticipatrice di nuove metodologie didattiche.

ALGEBRA ASTRATTA E INSEGNAMENTO SECONDARIO

2. La diffusione delle idee algebriche fondamentali stentava, però, a decollare, nonostante gli sforzi di associazioni come la *Società di Scienze Matematiche e Fisiche Mathesis*, ben radicata sul territorio nazionale (RIZZI [6]). Un contributo determinante in tale direzione è stato quello di ZAPPA-PERMUTTI [8]. Infatti le nozioni di gruppo, di complesso in un gruppo, di sottogruppo e le questioni sui laterali intimamente connesse con le congruenze modulate su un sottogruppo, collegate al teorema di *Lagrange*, sono oggi oggetto di studio in molti corsi secondari. Le classi di resti di numeri interi hanno perduto molto del loro, per così dire, misterioso fascino, per entrare in un ruolo di normale conoscenza. Il livello matematico generale si è molto innalzato, nonostante alcuni ambienti scolastici tradizionalmente non innovatori, in particolare i licei scientifici, abbiano, in un certo senso, remato contro.

Il concetto di gruppo, e più in particolare di gruppo di trasformazioni, ha contribuito a mettere a fuoco l'importanza dell'idea di *invariante*; idea che in qualche misura permette di superare il classico dualismo tra assoluto e relativo (SPERANZA [7]).

Vale la pena di ricordare anche che nel libro scolastico *La matematica, strutture, I*, Einaudi Scuola, Milano 1995, di L. Citrini, E. Castagnola e M. Impedovo si trova ancora oggi lo spettro delle nuove idee sulla didattica della matematica prodotte dalle originarie proposte di Frascati.

SUL CONCETTO DI IPEROPERAZIONE

3. Allo stato attuale, per proseguire in una graduale algebrizzazione della geometria, ha un certo interesse didattico la nozione di *iperoperazione*, che è

alla base di avanzate ricerche su *ipergruppi* e *join-space* (CORSINI [2]). Senza voler approfondire, nel prosieguo, questioni che esulano dal livello elementare di questa trattazione, ricordiamo che, dato un insieme H , e detto $P'(H)$ l'insieme delle sue parti non vuote, una iperoperazione $\bullet$ è un'operazione che associa a due *punti* di H una e una sola parte non vuota di H , ovvero uno ed un solo elemento di $P'(H)$, nel seguente modo:

$$(a,b) \in H^2 \rightarrow a \bullet b \in P'(H). \quad (1)$$

La (1) sintetizza il fatto che l'iperoperazione è un'applicazione φ tra il prodotto cartesiano $H^2 = H \times H$ e l'insieme delle parti non vuote di H stesso,

$$\varphi : H^2 \rightarrow P'(H).$$

L'operazione può essere reiterata come indicato in (2) e (3).

$$(a \bullet b) \bullet c = \bigcup_{z \in a \bullet b} z \bullet c \quad (2)$$

$$a \bullet (b \bullet c) = \bigcup_{w \in b \bullet c} a \bullet w \quad (3)$$

UN ESEMPIO DIDATTICO: UNA IPEROPERAZIONE SU UN CERCHIO

4. Ci si domanda ora se sia possibile introdurre il concetto di iperoperazione nella geometria elementare, per potere interpretare i fatti geometrici con un linguaggio puramente algebrico. La risposta può essere affermativa se riusciamo ad entrare in uno spirito di ricerca didattica, senza coltivare pregiudizi *antifusionisti* (cfr. § 2).

Dato, per esempio, un cerchio H di centro O e raggio qualsiasi, sia L la circonferenza ad esso associata, che consideriamo inclusa nel cerchio stesso. L'insieme $P'(H)$ delle parti non vuote di H contiene la stessa circonferenza, archi e corde comunque costruiti, segmenti e settori circolari, cerchi interni e figure continue e discrete immerse in H e costruite a piacere.

4.1. L'iperoperazione che prendiamo in considerazione sia definita nel modo seguente.

- (a) Ad ogni coppia di punti distinti del cerchio H risulti associato il segmento $P_1 P_2$ che li unisce, con l'ulteriore convenzione che, se i punti sono coincidenti, il risultato è il punto stesso pensato come *singleton* di $P'(H)$.

$$\text{Avremo} \quad P_1 \bullet P_2 = P_1 P_2 \quad ; \quad P_1 \bullet P_1 = P_1.$$

- (b) Ottenuto il segmento $P_1 \bullet P_2$ ¹ si costruisca l'unione di tutti i segmenti aventi come estremo fisso $P_3 \in H$ e come secondo estremo un punto M_K del segmento $P_1 \bullet P_2$.

$$\text{Avremo } (P_1 \bullet P_2) \bullet P_3 = \bigcup_{M_K \in P_1 \bullet P_2} M_K \bullet P_3 . \quad (\text{cfr. (2)})$$

- (c) Ottenuto il segmento $P_2 \bullet P_3$ si costruisca l'unione di tutti i segmenti aventi come estremo fisso P_1 e come secondo estremo un punto N_K del segmento $P_2 \bullet P_3$.

$$\text{Avremo } P_1 \bullet (P_2 \bullet P_3) = \bigcup_{N_K \in P_2 \bullet P_3} P_1 \bullet N_K . \quad (\text{cfr. (3)})$$

4.2. E' immediato verificare che:

- l'iperoperazione introdotta è commutativa;
- l'iperoperazione è associativa; ad ogni terna di punti corrisponde il loro triangolo se essi non sono allineati ed il segmento che li contiene se sono allineati;
- ad ogni coppia di punti distinti appartenenti alla circonferenza L corrisponde una corda del cerchio;
- ad ogni terna di punti appartenenti alla circonferenza corrisponde un triangolo inscritto in essa;
- allargando ulteriormente l'iperoperazione ad una quaterna di punti scelti sulla circonferenza, si osserva la costruzione di un quadrilatero inscritto nella circonferenza. Se si continua ad aumentare il numero dei punti si ottiene una successione di poligoni inscritti che "ammettono come limite" il cerchio stesso.

4.3. E' possibile inoltre, analogamente a quanto fatto in 4.1 e con le opportune ovvie precisazioni, definire algebricamente le figure immerse nel cerchio H di centro O nel seguente modo.

Segmento circolare; $\bigcup_{X \in P_1 P_2} (P_1 \bullet P_2) \bullet X$; $P_1 P_2$: arco di L .

Semicerchio; $\bigcup_{X \in P_1 P_2} (P_1 \bullet P_2) \bullet X$; $O \in P_1 \bullet P_2$; $P_1 P_2$: arco di L .

¹ Per semplicità si suppone che, nel seguito, i punti comunque considerati siano distinti.

Settore circolare;

$$\left(\bigcup_{\substack{Y \in P_1 \bullet P_2 \\ P_1, P_2 \in L}} 0 \bullet Y \right) \cup \left(\bigcup_{X \in P_1 P_2} (P_1 \bullet P_2) \bullet X \right).$$

Cerchio come luogo di corde;

$$\bigcup_{P_k, P_i \in L} P_k \bullet P_i = H.$$

CONCLUSIONI

5. Concludiamo questa breve nota osservando che l'iperoperazione introdotta nel cerchio H può rigenerare H stesso in più modi, ma comunque sempre con procedimenti iterativi. L'esempio fatto nel paragrafo precedente fornisce un'idea del campo di applicazione didattica del concetto di iperstruttura e, in particolare, di iperoperazione.

In tal senso quindi l'introduzione di particolari iperoperazioni potrà fornire validi *strumenti didattici* a sussidio dell'insegnamento (e quindi dell'apprendimento) della geometria piana e spaziale. Consentirà anche la fusione, nel senso richiamato al § 1, con i concetti di corrispondenza, applicazione e funzione, nonché delle loro applicazioni a *modelli matematici* della realtà, particolarmente in campo fisico, biologico ed economico.

BIBLIOGRAFIA

1. L. BERARDI, F. EUGENI e S. INNAMORATI, *Generalized designs, linear spaces, hypergroupoids and algebraic cryptography* - Proceedings of the Fourth International Congress on Algebraic Hyperstructures and Applications, Xanthi, Greece, World Scientific, (1990), 55-65.
2. P. CORSINI, *Recenti risultati in teoria degli ipergruppi*, B.U.M.I., 2A (1983), 135-138.
3. P. CORSINI, *Prolegomena of Hypergroup theory*, Aviani Editore, Udine 1995.
4. C. D'ANIELLO, *Bruno Rizzi e il fusionismo*, su questo stesso volume, 31-34.
5. E. GELSOMINI, *Un ipergruppo associato all'insieme delle matrici quadrate di ordine n non singolari*, su questo stesso volume, 99-103.
6. B. RIZZI, *de Finetti e il Periodico di Matematiche*, Per. Mat., 2/3 (1995), 69-76.
7. F. SPERANZA, *Aspetti matematici e fisici della epistemologia della Geometria*, Atti del XVII Convegno UMI, Edizioni della Unione Matematica Italiana, Bologna, (1995), 29-35.
8. G. ZAPPA e R. PERMUTTI, *Gruppi, corpi, equazioni*, Feltrinelli, Milano 1972.

UN NUOVO APPROCCIO ALLA TEORIA DEI SISTEMI DIGITALI MULTICANALE

Fabio Mercanti*, Luigi Cerritelli*, Vincenzo Di Marcello**

SUNTO - Si propone l'uso di un nuovo *formalismo* per lo studio della teoria dei *Sistemi Digitali Multicanale* (SDM).

ABSTRACT - A new formalism to study of the theory of the *Multichannel Digital Systems* (SDM) has been formulated.

INTRODUZIONE

0. Lo studio della teoria dei *Sistemi Digitali Multicanale* (SDM), generalizzazione del concetto di *sistema* (vedere per esempio BERTALAMFFY [1]) inteso nella sua più ampia accezione, presenta non poche difficoltà dovute al formalismo di cui la teoria stessa fa uso. Scopo di questo lavoro è quello di introdurre un formalismo alternativo, sperabilmente più conveniente, che rientri in quello *tradizionale della matematica*.

Si ricorda rapidamente che un SDM (MELZI-MERCANTI [6], MELZI [4]) è un sistema *dinamico* (RINALDI [9]) per il quale lo *stato iniziale*, l'*input-output* e lo *stato istantaneo* sono costituiti da informazione codificata

* Politecnico di Milano, via Bonardi 3.

** Mathesis di Teramo.

secondo le regole di un'algebra $\mathcal{A}$, sulle quali si basa un certo meccanismo *algebrico-combinatorio*, per il quale si rimanda a MELZI-MERCANTI [5], MELZI [3].

Con opportune precisazioni, ivi si considera un *alfabeto* T costituito da un numero finito di caratteri. Le *parole* su T costituiscono una *grammatica* G . Detto $P(G)$ l'insieme delle parti di G , gli elementi di $P(G)$ formano un'algebra di Boole $\mathcal{B}$. Introducendo alcune altre opportune *relazioni*, oltre a quelle sussistenti in $\mathcal{B}$, e similmente altre *operazioni*, dette *$\mathcal{A}$ -operazioni*, tra le parole di G , si ottiene una *struttura algebrica* $\mathcal{A}$, estensione di $\mathcal{B}$. Le *$\mathcal{A}$ -operazioni* (come suggerisce anche il loro nome: *selezione*, *selezione destra*, *selezione sinistra*, *taglio e concatenazione*) sono *verosimili modelli* delle operazioni di *sintesi polipeptidica*.

Il fatto che un SDM sia descrivibile con particolari meccanismi algebrico-combinatori, agenti su informazione opportunamente codificata, lo rende particolarmente adatto alla formulazione di *modelli matematici* per lo studio di vari oggetti e fenomeni naturali, come ad esempio *sistemi biologici*, *sistemi economici*, *sistemi di edifici rispetto ai processi di manutenzione e degrado*, *fenomeni sismici*, ecc.. Se infatti l'idea tradizionale di sistema è dominata dai concetti di *proporzionalità e continuità*, l'idea di SDM è invece dominata dalla nozione di *soglia*, ossia dal fenomeno per il quale stimoli inferiori ad una certa intensità critica non provocano alcuna risposta in un dato sistema (MERCANTI [8]). Un esempio tipico di risultato in tal senso può trovarsi in MERCANTI [7], dove si *descrive* e si *giustifica* un particolare fenomeno della percezione, ritrovando e giustificando i risultati che si ottengono *sperimentalmente*. L'impostazione assiomatica delle teorie di un SDM forma una base interpretativa dei fenomeni nervosi, nel senso che consente di fondare una *teoria matematica dei fenomeni mentali* in senso lato.

Un SDM è la naturale generalizzazione del concetto di *neuromacchina* (MELZI [2])¹, dalla quale ha ereditato il formalismo che in questa trattazione si vuole rendere più semplice.

¹ Uno stadio, per la verità, intermedio tra il concetto di neuromacchina e quello di SDM è costituito da un *semiautoma* (MELZI [3]), non strettamente necessario alla presente trattazione.

PROCESSI DI ELABORAZIONE DI UNA NEUROMACCHINA

1. Una *neuromacchina* $\mathcal{N}$ è costituita da un insieme di *moduli* comunicanti tra loro in modo opportuno. In sintesi, il funzionamento di $\mathcal{N}$ è così descrivibile.

Siano $I_{\mathcal{N}}$, $O_{\mathcal{N}}$ rispettivamente gli insiemi degli *stimoli* e delle *risposte* di $\mathcal{N}$. Gli elementi di $I_{\mathcal{N}}$, $O_{\mathcal{N}}$ sono *parole*², cioè *stringhe* costruite secondo una certa *grammatica* di un dato *alfabeto*, ai cui simboli (neurocaratteri) si aggiungono i simboli metalinguistici « , », « * », costruendo in tal modo certe *neuroparole*. Le neuroparole sono, quindi, stringhe costruite secondo la grammatica, su un neuroalfabeto comprendente un numero finito di simboli.

Ad ogni *stimolo* corrisponde una *risposta* di $\mathcal{N}$. In altri termini, una volta che sia stata introdotta nella neuromacchina $\mathcal{N}$ una parola $v \in I_{\mathcal{N}}$ la $\mathcal{N}$ proclama, entro un intervallo di tempo finito, una nuova parola $w \in O_{\mathcal{N}}$. La risposta di $\mathcal{N}$ è funzione dello stimolo v in entrata, nonché dello stato iniziale di $\mathcal{N}$ (stato α). Il tempo è suddiviso in intervalli che, per comodità, sono supposti di uguale durata, suddivisi in intervalli *dispari* ed in intervalli *pari*.

² Nel modulo risiedono *caratteri materiali*, copie materiali dei caratteri di un *alfabeto finito* $c_0, c_1, \dots, c_{h-1}$, detto *alfabeto materiale*, comprendente un simbolo privilegiato denotato con « • », in generale coincidente con c_0 . Una neuroparola materiale è un allineamento ordinato dei precedenti neurocaratteri. In particolare, si chiama *liscia* una neuroparola non contenente occorrenze del carattere « • » e si usano le lettere latine maiuscole nel senso di variabili, indicanti neuroparole materiali non ridotte a singoli neurocaratteri. Inoltre, con evidente significato dei simboli, si dice che neuroparole del tipo $\bullet A, B\bullet, \bullet C\bullet$ sono, rispettivamente, *chiusa a sinistra*, *chiusa a destra* o, semplicemente, *chiusa*. Analogamente si può parlare di neuroparole *aperte* nello stesso senso. Le neuroparole lisce, con cui terminano le neuroparole aperte a destra, prendono il nome di *suffissi*; le neuroparole, con cui iniziano le neuroparole aperte a sinistra, sono dette *prefissi*. Per esempio, nelle neuroparole materiali $\bullet A_1 \bullet A_2 \bullet \dots \bullet A_r \bullet B, \bullet C \bullet D_1 \bullet \dots \bullet D_s \bullet E, F \bullet G_1 \bullet \dots \bullet G_t, B, E$ sono suffissi e C, F prefissi. Si dice, inoltre, che una neuroparola materiale è *feconda*, se è composta da neuroparole materiali lisce tutte chiuse, tranne eventualmente l'ultima, che può essere aperta a destra. Per esempio le neuroparole $A_1 \bullet A_2 \bullet \dots \bullet A_r \bullet$ e $\bullet A_1 \bullet A_2 \bullet \dots \bullet A_r \bullet B$, con $A_1, A_2, \dots, A_r, B$ tutte lisce, sono entrambe feconde. Si conviene, infine, che l'insieme delle neuroparole materiali comprenda una neuroparola materiale *muta*. La costruzione delle neuroparole avviene, con evidente significato dei termini, a secondo che le parole stesse siano o no aperte, aperte a destra, aperte a sinistra (MELZI [2]).

Il meccanismo secondo cui avviene l'elaborazione entro la neuromacchina è stabilito dagli assiomi di MELZI [2].³

PUNTO DI VISTA ALTERNATIVO

2. Vogliamo ora considerare il funzionamento di una neuromacchina da un punto di vista leggermente differente. Negli intervalli di tempo (cfr. § 1) dispari⁴ l'informazione può entrare oppure uscire da una neuromacchina sotto forma di una sola neuroparola o di più neuroparole separate fra loro dal simbolo « , », cioè sotto forma di una parola non contenente il simbolo « * ».

Trascuriamo gli intervalli di tempo pari e introduciamo la variabile temporale t assumente valori interi (non negativi, se l'origine dei tempi è scelta coincidente con l'inizio dell'elaborazione).

Si può, così, descrivere un generico processo di elaborazione da parte di una neuromacchina dicendo che in *ogni* intervallo di tempo t entra in $\mathcal{N}$ un'informazione costituita da una parola x_t ed esce da $\mathcal{N}$ un'informazione costituita da una parola y_t . Tali parole x_t, y_t non contengono il simbolo « * »

³ PRIMO ASSIOMA. Ogni modulo $\mathcal{H}_j$ componente una neuromacchina $\mathcal{H} = \mathcal{H}_j$ ($j = 1, 2, \dots, n, \dots$) contiene inizialmente un insieme α_j , funzione di j , di neuroparole materiali allo stato latente, e ricostruisce tutte e sole queste neuroparole durante ogni intervallo temporale δ_i di indice i pari quando esse sono state cancellate nell'intervallo precedente δ_{i-1} .

SECONDO ASSIOMA. Le neuroparole binarie circolano in $\mathcal{H}_j$ solo negli intervalli temporali δ_i di indice i dispari. Se in un intervallo temporale δ_i di indice dispari circola in $\mathcal{H}_j$ una neuroparola binaria $\gamma^{-1}(XY)$, corrispondente di una neuroparola materiale XY , e in $\mathcal{H}_j$ risiedono una neuroparola materiale del tipo $\bullet A \bullet X$ e una neuroparola materiale del tipo $Y \bullet B$ (X e Y lisce e aperte, A aperta, B aperta a sinistra, lisce o no, eventualmente vuote), allora entro l'intervallo temporale δ_i le neuroparole materiali $\bullet A \bullet X$ e $Y \bullet B$ vengono cancellate da $\mathcal{H}_j$ e in questo modulo vengono a risiedere le neuroparole materiali $\bullet A \bullet$ e $\bullet B$.

TERZO ASSIOMA. Se in un intervallo temporale δ_i di indice i dispari nel modulo $\mathcal{H}_j$ si forma una neuroparola materiale chiusa del tipo $\bullet A_1 \bullet A_2 \bullet A \bullet \dots \bullet A_r \bullet$, con $A_1, A_2, \dots, A_r$ lisce, allora tale neuroparola viene cancellata nel corso dell'intervallo pari successivo δ_{i+1} e nell'intervallo dispari δ_{i+2} in $\mathcal{H}_j$ circolano le copie binarie delle neuroparole A_h ($h = 1, 2, \dots, r$).

QUARTO ASSIOMA. Se in un intervallo temporale δ_i di indice i dispari in un modulo $\mathcal{H}_j$ circola la copia binaria $\gamma^{-1}(A_h)$ di una componente liscia di una neuroparola materiale neoformata in $\mathcal{H}_j$, la neuroparola binaria $\gamma^{-1}(A_h)$ circola, nello stesso intervallo di tempo, in tutti i moduli di $\mathcal{H}$ (MELZI [2]).

⁴ Gli intervalli di tempo pari servono unicamente per ricostituire lo stato iniziale α della neuromacchina (cfr. § 1).

e possono essere mute. In generale le parole x_t sono formate da neuroparole costituenti in parte *dati* da elaborare e in parte *risultati parziali* dell'elaborazione⁵. In tal senso ogni y_t è funzione delle parole $x_{t'}$ ($t' < t$).

MATRICI COME RAPPRESENTAZIONI DI PAROLE

3. Siano $\mathcal{A}_{\mathcal{N}}$, $\mathcal{B}_{\mathcal{N}}$ rispettivamente, gli insiemi delle neuroparole con cui sono formate le parole *stimoli* e *risposte* di $\mathcal{N}$, cioè gli elementi di $I_{\mathcal{N}}$, $O_{\mathcal{N}}$ (cfr. § 1). Consideriamo l'insieme *unione* di $\mathcal{A}_{\mathcal{N}}$, $\mathcal{B}_{\mathcal{N}}$, cioè l'insieme

$$\mathcal{U}_{\mathcal{N}} = \mathcal{A}_{\mathcal{N}} \cup \mathcal{B}_{\mathcal{N}}.$$

Supponiamo $\mathcal{U}_{\mathcal{N}}$ ordinato secondo una legge qualsiasi e denotiamo con U_h ($h = 1, 2, \dots, m$) gli elementi di $\mathcal{U}_{\mathcal{N}}$ ordinati secondo tale legge. Ogni U_h può, così, essere posto in corrispondenza biunivoca con una matrice unicolonnare u_h a n righe, stabilendo che a U_h corrisponda la matrice i cui elementi sono tutti nulli, ad eccezione di quello della h -esima riga che è uguale a 1. Si pone, cioè

$$U_h \leftrightarrow u_h = \begin{bmatrix} 0 \\ \dots \\ 0 \\ 1 \\ 0 \\ \dots \\ \dots \\ 0 \end{bmatrix}. \quad (1)$$

In tale modo ogni neuroparola appartenente ad $\mathcal{A}_{\mathcal{N}}$, oppure a $\mathcal{B}_{\mathcal{N}}$, può essere individuata mediante una matrice unicolonnare u_h .

Analogamente ad ogni parola, costituita da neuroparole separate fra loro dal simbolo « , », si può associare una matrice unicolonnare, a n righe, i cui elementi sono tutti nulli, ad eccezione di quelli, uguali a 1, corrispondenti alle neuroparole che formano la parola in questione. Si stabilisce, cioè, la corrispondenza

⁵ Si può supporre con MELZI [2] che un *osservatore* Ω possa opportunamente interagire con $\mathcal{N}$, proclamando certe neuroparole. In tal senso, le parole x_t sono costituite da neuroparole proclamate sia da Ω (*dati da elaborare*) sia da $\mathcal{N}$ (*risultati parziali*).

$$V = U_i, U_j, \dots, U_l \leftrightarrow v = \begin{bmatrix} 0 \\ \dots \\ 1 \\ \dots \\ 1 \\ \dots \\ 1 \\ \dots \\ \dots \\ 0 \end{bmatrix}. \quad (2)$$

Con tale convenzione la parola muta è descritta da una matrice i cui elementi sono tutti nulli.

4. Se in uno stesso intervallo di tempo circolano le parole (cfr. (2))

$$V = U_i, U_j, \dots, U_l,$$

$$W = U_a, U_b, \dots, U_c,$$

si può equivalentemente dire che circola l'unica parola

$$Z = V, W = U_i, U_j, \dots, U_l, U_a, U_b, \dots, U_c.$$

Se v, w, z sono le matrici (1) associate, rispettivamente, alle parole V, W, Z , vale la relazione

$$z = v \oplus w,$$

nella quale l'operazione indicata con il simbolo « $\oplus$ » è una particolare *somma*. Essa è analoga all'usuale somma matriciale, purché valgano le seguenti regole

$$\begin{cases} 0 \oplus 0 = 0 \\ 0 \oplus 1 = 1 \oplus 0 = 1 \\ 1 \oplus 1 = 1 \end{cases}$$

L'operazione inversa non viene definita.

EQUAZIONI CANONICHE DELLE NEUROMACCHINE

5. Siano x_t, y_t le matrici (2) corrispondenti rispettivamente alle parole X_t, Y_t . Il processo di elaborazione dei dati mediante una neuromacchina può essere considerato, così, come il calcolo delle y_t in funzione delle x_t , essendo $t' < t$ (cfr. § 2). Nel caso di un processo comprendente $n+1$ passi di elaborazione (cioè di durata uguale a $n+1$ intervalli di tempo) si possono, pertanto, scrivere le seguenti relazioni

$$\begin{cases} y_1 = f_1(x_0) \\ y_2 = f_2(x_0, x_1) \\ \dots\dots\dots \\ y_n = f_n(x_0, x_1, \dots, x_{n-1}), \end{cases} \quad (3)$$

dove $f_1, f_2, \dots, f_n$ denotano funzioni qualsiasi generalmente non lineari. Poiché le parole x_t in entrata di $\mathcal{N}$ sono costituite da neuroparole proclamate in parte dall'osservatore e in parte dalla stessa neuromacchina⁶, le x_h , con $h = 1, 2, \dots, n-1$, si possono scomporre come segue

$$x_h = y_h \oplus x'_h, \quad (4)$$

dove x'_h rappresenta il contributo dell'osservatore, o, in generale, dell'esterno.

Tenendo conto delle (4) le (3) possono, così, essere scritte nella *forma canonica*

$$\begin{cases} y_1 = F_1(x_0) \\ y_2 = F_2(x_0, x'_1, y_1) \\ y_3 = F_3(x_0, x'_1, x'_2, y_1, y_2) \\ \dots\dots\dots \\ y_n = F_n(x_0, x'_1, \dots, x'_{n-1}, y_1, y_2, \dots, y_{n-1}). \end{cases} \quad (5)$$

Poiché ogni risposta di ogni neuromacchina è proclamata entro un tempo finito, per ogni neuromacchina esiste un numero finito di funzioni, $F_1, F_2, \dots, F_n$ che ne caratterizzano il comportamento. Ogni neuromacchina

⁶ Cfr. Nota ⁵.

può, così, essere considerata equivalente ad una opportuna n -pla di funzioni $F_1, F_2, \dots, F_n$.

CONCLUSIONI

6. Si presentano i seguenti due problemi fondamentali.

Problema diretto. Dato lo stimolo, determinare la risposta, cioè, date le quantità $x_0, x'_1, \dots, x'_{n-1}$, determinare le quantità $y_1, y_2, \dots, y_n$.

Se sono note le funzioni $F_1, F_2, \dots, F_n$, ciò si ottiene immediatamente dalle (5).

Problema inverso. Data la risposta, determinare lo stimolo (o gli stimoli) che l'hanno causata, cioè, date le quantità $y_1, y_2, \dots, y_n$, determinare le quantità $x_0, x'_1, \dots, x'_{n-1}$.

Ciò richiede la risoluzione delle (5) rispetto alle $x_0, x'_1, \dots, x'_{n-1}$.

7. Le risposte ai problemi diretto e inverso del § 6, ottenute passando attraverso le formulazioni (1) e (2), forniscono una palese esemplificazione della semplificazione del linguaggio finora usato nello studio delle neuromacchine e, quindi, degli SDM. La semplificazione non è solo simbolica, ma consente anche l'uso di strumenti classici della matematica. In tal senso è attualmente allo studio una riformulazione dell'algebra $\mathcal{A}$ di [5], [3] (cfr. § 0) secondo le risultanze sopra esposte.

BIBLIOGRAFIA

1. L. von BERTALAMAFFY, *Teoria generale dei sistemi*, ISEDI, Milano 1971.
2. G. MELZI, *Una definizione assiomatica del concetto di neuromacchina*, Rend. Sem. Mat. di Brescia, 5 (1981), 9-74.
3. G. MELZI, *Sulla definizione di semiautoma*, Rend. Sem. Mat. di Brescia, 10 (1988), 23-50.
4. G. MELZI, *Sistemi dinamici digitali e loro possibili applicazioni*, Ratio Math, 2 (1991), 157-160.
5. G. MELZI e F. MERCANTI, *An algebraic-combinatory Theory of real nervous System*, Rend. Sem. Mat. di Brescia, 9 (1988), 107-121.
6. G. MELZI e F. MERCANTI, *Teoria dei sistemi e Sistemi Digitali Multicanale*, Sinopie, Politecnico di Milano, 2 (1989), 34-35.
7. F. MERCANTI, *Una suggestiva analogia fra fenomeni biologici e fenomeni economici a soglia*, Ratio Math., 2 (1991), 161-165.
8. F. MERCANTI, *La descrizione matematica dei fenomeni a soglia*, Per.di mat., Vol. 69, 2 (1993), 56-60.
9. S. RINALDI, *Teoria dei sistemi*, Clup., Politecnico di Milano 1977.

UN IPERGRUPPO ASSOCIATO ALL' INSIEME DELLE MATRICI QUADRATE DI ORDINE N NON SINGOLARI

Emma Gelsomini*

SUNTO - Si assegna un' *iperoperazione* sull'insieme delle matrici quadrate non singolari di ordine n , che conferisce all'insieme una struttura di *ipergruppo non commutativo*.

ABSTRACT - An *hyperoperation* on the set of the not singular square matrices of the n -order has been assigned. This hyperoperation induces in the set a structure of *not commutative hypergroup*.

0. In questo lavoro si considera l'insieme M delle *matrici quadrate* di ordine n , non singolari, e una operazione σ definita ponendo

$$\forall x, y \in M \quad x \sigma y = \{ k \ x \ y \}_{k \in \mathfrak{N}}, \quad (1)$$

essendo xy l'ordinario prodotto righe per colonne tra matrici. Si dimostra che la struttura

$$(M, \sigma) \quad (2)$$

è un *ipergruppo*, non commutativo.

* Lavoro svolto nell'ambito di un gruppo di ricerca MPI 60 % del Politecnico di Milano.

1. Si ricordano brevemente, nel seguito, i concetti di *ipergruppoide* (1.1), di proprietà *associativa* (1.2) e di proprietà di *riproducibilità* (1.3), per un ipergruppoide (CORSINI [2], [3]). Le proprietà 1.1, 1.2 e 1.3 individuano un ipergruppo e dovranno, evidentemente, valere per la coppia (M, σ) (cfr.(2)).

1.1 Un ipergruppoide è una coppia $(G, \circ)$, dove G è un insieme non vuoto e

$$\circ : G \times G \rightarrow P'(G) \quad (3)$$

è una *applicazione* di $G \times G$ nell'insieme $P'(G)$ delle parti non vuote di G .

1.2 Un ipergruppoide che goda della proprietà associativa

$$\forall x, y, z \in G \quad (x \circ y) \circ z = x \circ (y \circ z), \quad (4)$$

dove

$$(x \circ y) \circ z = \bigcup_{a \in x \circ y} a \circ z$$

e

$$x \circ (y \circ z) = \bigcup_{b \in y \circ z} x \circ b,$$

viene detto *semipergruppo*.

1.3 Un semipergruppo soddisfacente la seguente proprietà

$$\forall a, b \in G \quad \exists x \in G : b \in a \circ x, \exists y \in G : b \in y \circ a, \quad (5)$$

detta *proprietà di reproducibilità*, è un *ipergruppo*.

2. Si osservi dapprima che per (M, σ) (cfr.(2)) vale la proprietà 1.1, ovvero che (cfr. (3))

$$\sigma : M \times M \rightarrow P'(M),$$

è, manifestamente, un'applicazione.

3. Si verifica ora che per (M, σ) vale anche la proprietà 1.2., cioè si dimostra che l'ipergruppoide appena definito nel § 2 gode della proprietà associativa (cfr. 1.2).

3.1. Si considerino infatti tre matrici $m, n, p \in M$. La (4) del § 1.2 che regola la proprietà associativa, diventa in questo caso

$$\forall m, n, p \in M, \quad (m \sigma n) \sigma p = m \sigma (n \sigma p). \quad (6)$$

Con evidente significato dei simboli e delle formule (cfr. (1)), si ottiene, successivamente,

$$\begin{aligned} (m \sigma n) \sigma p &= \{kmn\}_{k \in \mathfrak{R}} \sigma p = \\ &= \{kh(mn)p\}_{k, h \in \mathfrak{R}} = \{t(mn)p\}_{t \in \mathfrak{R}}, \end{aligned} \quad (7)$$

con $t = kh$.

Analogamente si ha:

$$\begin{aligned} m \sigma (n \sigma p) &= m \sigma \{h'np\}_{h' \in \mathfrak{R}} = \\ &= \{h'k'm(np)\}_{h', k' \in \mathfrak{R}} = \{t'm(np)\}_{t' \in \mathfrak{R}}, \end{aligned} \quad (8)$$

con $t' = h'k'$.

Tenendo conto della proprietà associativa del prodotto tra matrici, il confronto tra la (7) e la (8) consente l'immediata verifica della (6).

4. Si verifica infine che per (M, σ) vale la proprietà 1.3, ovvero si dimostra che il *semipergruppo* appena definito (cfr. 1.2, (7) e (8)) gode della proprietà di riproducibilità (5)

$$\begin{aligned} \forall m, n \in M, \quad \exists x \in M: n \in m \sigma x &= \{kmx\}_{k \in \mathfrak{R}}, \\ \exists y \in M: n \in y \sigma m &= \{kym\}_{k \in \mathfrak{R}}, \end{aligned} \quad (9)$$

essendo m e n matrici non singolari.

4.1. In entrambi i casi si ponga

$$\begin{aligned}x &= m^{-1}n, \\ y &= n m^{-1},\end{aligned}$$

essendo l'esistenza di m^{-1} garantita dalla non singolarità di $m \in M$.

4.2. Con evidente significato di linguaggio e simboli, si calcola immediatamente

$$m \sigma x = \{k m x\}_{k \in \mathfrak{M}} = \{k m m^{-1} n\}_{k \in \mathfrak{M}} = \{k n\}_{k \in \mathfrak{M}}.$$

Nel caso in cui $k=1$, è chiaramente verificato che

$$n \in m \sigma x.$$

Analogamente si può concludere anche che

$$n \in y \sigma m.$$

La (9) è quindi verificata.

4.3. Pertanto si conclude che il semipergruppo considerato gode anche della proprietà di riproducibilità (cfr. 1.3).

5. Da quanto detto nei paragrafi 2, 3 e 4, segue che la coppia (M, σ) è un ipergruppo.

BIBLIOGRAFIA

1. L. BERARDI, F. EUGENI e S. INNAMORATI, *Generalized designs, linear spaces, hypergroupoids and algebraic cryptography*, Proceedings of the Fourth International Congress on Algebraic Hyperstructures and Applications, Xanthi, Greece, World Scientific, (1990), 55-65.
2. P. CORSINI, *Recenti risultati in teoria degli ipergruppi*, B.U.M.I., 2-A (1983), 133-138.
3. P. CORSINI, *Prolegomena of hypergroup theory*, Aviani Editore, Udine 1992.
4. M.GIONFRIDDO, *Hypergroups associated with multihomomorphisms between generalized graphs*, Atti Convegno su "Sistemi binari e loro applicazioni", Taormina (1978), 161-174.