

Quaderni di informatica

4

Rassegna di tecnologia ed applicazioni degli elaboratori elettronici - Anno II - Numero 2 - 1975



Honeywell
Honeywell Information Systems Italia

FPDM-31

Quaderni di informatica

Rassegna di tecnologia ed applicazioni degli elaboratori elettronici
Anno II - Numero 2 - 1975

Sommario

Le strutture informative

I dati che intervengono in una elaborazione possono essere inquadrati concettualmente secondo « strutture informative astratte ». L'articolo espone i principi di questa schematizzazione logica che costituisce un aspetto fondamentale dell'informatica e rappresenta la premessa per la organizzazione effettiva dei dati nella memoria dell'elaboratore. A questo secondo aspetto, che costituisce le « strutture informative concrete », verrà dedicato un successivo articolo della rivista.

M. ITALIANI

Problemi nella interazione uomo-calcolatore

L'interfaccia tra l'uomo ed il calcolatore costituisce notoriamente un punto cruciale nell'impiego di questo mezzo, in particolare allorché l'operatore non è uno specialista di calcolatori, ma un normale utente. L'articolo focalizza questa problematica attraverso una indagine condotta su un'ampia casistica di utenti appartenenti a diverse categorie professionali.

K.E. EASON, L. DAMORADAN, T.F. STEWART

Metodi automatici di gestione delle scorte nelle aziende di distribuzione

La gestione del magazzino è un problema classico della automazione aziendale. Però, mentre in passato ci si limitava alla contabilità delle entrate e delle uscite, ora l'obiettivo si è allargato alla gestione automatica nel senso più ampio. Attraverso modelli di previsione delle vendite, l'elaboratore provvede oggi all'aggiornamento tempestivo delle scorte, emettendo automaticamente gli ordini di rimpiazzo.

R. BRESCHI

Che cosa possono fare e che cosa non possono fare le macchine

« Può una macchina pensare? ». « Può una macchina riprodurre se stessa? ». Queste domande, che l'uomo si è posto anche in altre epoche tecnologiche, hanno ricevuto nuova luce e accresciuto interesse dall'avvento del calcolatore, subito chiamato, antropomorficamente, « cervello elettronico ». A questi stimolanti interrogativi danno qui la loro risposta due noti scienziati.

J. NIEVERGELT, J.C. FARRAR

Rilievi sul mercato dei calcolatori in Italia

Questo articolo presenta una serie di informazioni statistiche sul parco degli elaboratori elettronici in Italia, cercando di mettere in relazione la dimensione e la struttura del parco con quella dei principali aggregati economici e, facendo, ove possibile, confronti con altri paesi. Questo studio vuole costituire, stante la carenza nel nostro paese di rilevazioni pubbliche, un punto di riferimento della situazione italiana e verrà aggiornato su queste pagine con periodicità annuale.

R. LEVRERO

Resoconti e recensioni

Direttore Responsabile
F. Filippazzi

Comitato di Redazione
A. Cicu, M. De Marco, S. Fabris,
C. Falcetti, R. Levrero, P. Lupo,
N. Minnaja, G. Occhini, G. Rapelli,
F. Ricci, M. Rossi

Redazione
Honeywell Information Systems Italia
Centro di Ricerca e Progettazione
20010 Pregnana Milanese (Milano)

Stampa
grafiche A. Nava, Milano

Autorizzazione del Tribunale di Milano
n. 93 del 20 Marzo 1974

«Quaderni di Informatica» ospita
articoli e contributi di varia provenienza.
La responsabilità delle opinioni
espresse rimane agli Autori.

La riproduzione di articoli
della rivista, o di parte di essi,
è consentita solo citando la fonte
e l'autore.

Le strutture informative

MARIO ITALIANI
*Consiglio Nazionale delle Ricerche
Laboratorio di Analisi Numerica
Pavia*

1. Introduzione

Un programma progettato per risolvere uno specifico problema in uno qualsiasi degli innumerevoli campi di applicazione dei calcolatori elettronici difficilmente si trova a dover trattare dati isolati, ma, nella quasi totalità dei casi, elabora dati appartenenti ad uno o più insiemi finiti.

L'appartenenza di un dato ad un determinato insieme è definita dall'esistenza di relazioni logiche più o meno complesse tra il dato stesso e gli altri dati dell'insieme. Nei casi più elementari tali relazioni logiche potranno consistere ad esempio semplicemente nel fatto che i dati si presentano in « input » al programma in un ordine determinato. Qualche volta invece le relazioni potranno essere meno semplici come nel caso di dati che rappresentano elementi di una matrice in un problema matematico o elementi di un reticolo PERT.

Questa strutturazione logica degli insiemi di dati, che è naturalmente indipendente dal modo in cui i dati sono rappresentati nella memoria del calcolatore, è stata studiata con sistematicità soltanto in tempi relativamente recenti, tanto che la nomenclatura in uso è tuttora piuttosto ambigua.

Al giorno d'oggi tuttavia essa costituisce un capitolo fondamentale per chi si inizia agli studi di Informatica e pertanto può essere interessante, anche in questa sede, passare in rassegna i principali tipi di insiemi di dati dotati di una struttura logica, cioè le cosiddette « strutture astratte di dati » o « strutture informative astratte ».

È conveniente poi che questa rassegna sia completata da un esame dei metodi più in uso per la rappresentazione delle strutture astratte nella memoria di un calcolatore, cioè dalle cosiddette « strutture concrete di memoria » o « strutture informative concrete ».

A queste sarà dedicato un secondo articolo che comparirà in un prossimo numero della rivista. Per quanto riguarda la nomenclatura che, come si è detto, è spesso discordante, si adotterà per le varie strutture un termine di riferimento, citando tuttavia anche altri eventuali termini comunemente utilizzati.

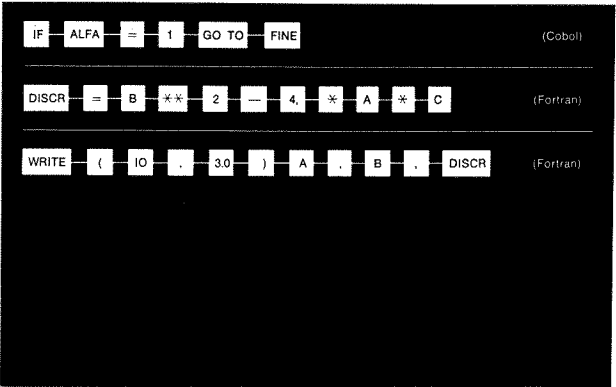
2. La «lista lineare»

Una struttura informativa astratta particolarmente semplice è la « lista lineare », termine con il quale si indica un insieme finito di dati, tra i quali è stabilito un ordine che permette di determinare quali

Fig. 1 - Esempio di lista lineare (Testo letterario).



Fig. 2 - Esempio di lista lineare (frasi di programmi simbolici).



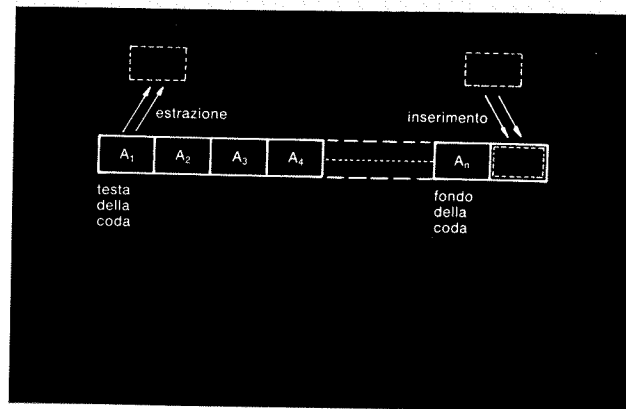


Fig. 3 - Rappresentazione schematica di una « coda ».

sono il primo e l'ultimo elemento. Inoltre l'accesso ad un particolare elemento dell'insieme avviene necessariamente per mezzo di una « ricerca sequenziale », cioè facendo scorrere gli elementi della lista, secondo l'ordine stabilito, a partire dal primo elemento.

Come esempio di lista lineare si può considerare un testo letterario (fig. 1) in cui gli elementi che costituiscono l'insieme sono le singole parole e i segni di interpunzione, mentre l'ordinamento è definito dalla posizione fisica che gli elementi occupano sul testo.

Se invece di un testo letterario si considerano ad esempio le frasi di un programma redatto in un linguaggio simbolico di programmazione come il FORTRAN o il COBOL (fig. 2), si ha ancora una lista lineare i cui elementi sono etichette, identificatori, costanti, parole chiave, operatori, parentesi, separatori, ecc.

Un particolare tipo di lista lineare, e precisamente quello in cui gli elementi componenti sono caratteri di un determinato alfabeto, viene normalmente chiamato « stringa », parola questa che è la traduzione un po' forzata del termine inglese « string » che in realtà ha in questa lingua un significato meno banale che non « laccio per le scarpe ».

Da qualche autore il termine stringa viene addirittura usato al posto del termine lista, tuttavia l'uso qui adottato sembra più soddisfacente.

È chiaro comunque che lo stesso esempio di testo letterario riportato della figura 1 può essere visto come una stringa i cui elementi sono caratteri dell'alfabeto usato (lettere maiuscole, lettere minuscole, segni di interpunzione, spazi, ecc.).

Gli elementi di una lista possono essere indicati con una notazione del tipo:

$$A_1, A_2, \dots, A_{i-1}, A_i, A_{i+1}, \dots, A_{n-1}, A_n$$

che mette in rilievo il loro ordinamento, anche se è opportuno ricordare che l'accesso ad un elemento delle liste lineari avviene con una ricerca sequenziale a partire da A_1 e non tramite il valore dell'indice. Il numero n dei dati che compongono una determinata lista può essere costante nel tempo, e si dirà allora che la lista è di « lunghezza fissa », oppure può variare nel tempo, e la lista si dice allora di « lunghezza variabile ».

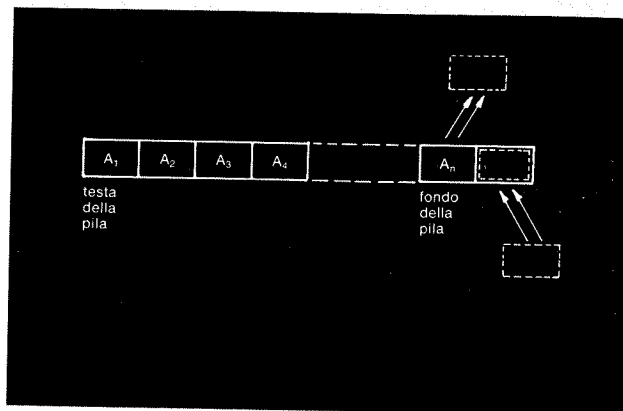


Fig. 4 - Rappresentazione schematica di una « pila ».

Le operazioni che si effettuano più frequentemente sugli elementi di una lista sono le seguenti:

- Accedere ad un elemento per esaminarlo o per modificarlo.
- Inserire un elemento prima o dopo un qualsiasi elemento A_i e in particolare inserire un elemento davanti al primo o dopo l'ultimo.
- Eliminare un elemento qualsiasi A_i e in particolare il primo elemento o l'ultimo.

Le ultime due operazioni sono ovviamente caratteristiche delle sole liste a lunghezza variabile; in particolare le liste a lunghezza variabile sulle quali le operazioni di estrazione avvengono solo sul primo o l'ultimo elemento e le operazioni di inserimento solo davanti al primo elemento o dopo l'ultimo prendono nomi particolari (« coda », « pila », « doppia coda ») in funzione del modo in cui tali operazioni vengono eseguite.

Prima di passare ad esaminare queste particolari liste è opportuno osservare che la nozione di lista lineare può essere generalizzata: si possono infatti considerare liste i cui elementi sono a loro volta liste e così via.

3. Coda, pila e doppia coda

Con il termine « coda » si indica una lista lineare di lunghezza variabile in cui tutti gli inserimenti avvengono dopo l'ultimo elemento (che prende il nome di « fondo » della coda), mentre tutte le estrazioni avvengono sul primo elemento (che prende il nome di « testa » della coda). Pertanto il primo elemento che può essere estratto in un determinato istante è quello, tra i presenti nella coda, che è stato inserito per primo. Per questa ragione alla coda viene dato anche il nome di lista « FIFO » dove questa sigla sta per « First-In-First-Out » (fig. 3).

Il termine coda ha un significato del tutto ovvio. Infatti la struttura ora descritta si adatta perfettamente a rappresentare le « code » o « file di attesa » che possono essere costituite da persone davanti ad uno sportello bancario (in questo caso ciascun elemento della coda potrà ad esempio comprendere un codice cliente, l'importo di un versamento, una casuale, ecc.) oppure da automobili in una officina

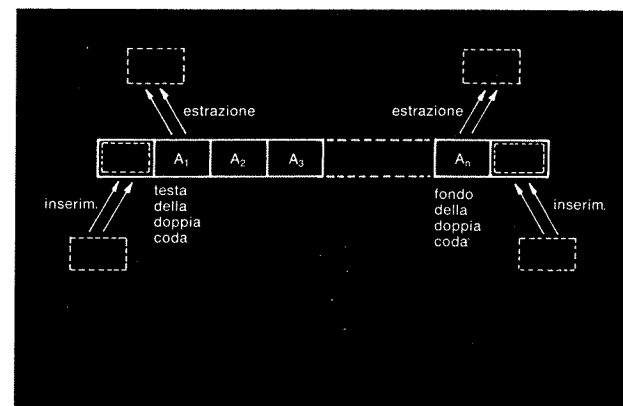


Fig. 5 - Rappresentazione schematica di una « doppia coda ».

riparazioni, da pezzi semilavorati davanti ad una macchina utensile e così via.

Il termine « pila », indica invece una lista lineare di lunghezza variabile in cui sia gli inserimenti che le estrazioni avvengono ad un solo estremo. Di conseguenza in una pila, che è nota anche con il termine « stack », il primo elemento che può essere estratto ad un determinato istante è quello che è stato inserito per ultimo: per questo motivo alla pila viene anche dato il nome di lista « LIFO » che sta per « Last-In-First-Out » (fig. 4).

Questo tipo di struttura trova applicazioni meno ovvie della precedente: essa trova impieghi di notevole rilievo nella realizzazione dei compilatori di linguaggi di programmazione e in particolare nel trattamento delle espressioni aritmetiche.

Il termine « doppia coda » indica infine una lista lineare di lunghezza variabile (di impiego poco frequente) in cui gli inserimenti e le estrazioni possono avvenire indifferentemente ad entrambi gli estremi come è indicato schematicamente nella fig. 5.

4. Matrici

Questo termine, che è stato utilizzato come corrispondente non del tutto appropriato del vocabolo anglosassone « array », il quale ha un significato più generale, indica un insieme finito di elementi in corrispondenza biunivoca con un insieme di n -ple ordinate di numeri interi detti « indici ». Se $n=2$ la matrice è detta a 2 dimensioni, mentre per $n=1$ si ha la matrice unidimensionale o vettore.

Nei casi più semplici (e più comuni) ciascun indice varia da un valore minimo ad un valore massimo con passo unitario come nella figura 6 in cui è schematizzata una matrice a due dimensioni con 4 righe e 5 colonne. Si può pensare che tale matrice rappresenti una tavola a doppia entrata e che il valore dell'elemento generico A_{ij} rappresenti la frequenza con cui due caratteristiche si presentano contemporaneamente in un gruppo di individui. Le due caratteristiche considerate sono, nell'esempio di fig. 6, l'altezza (suddivisa in 5 intervalli) e il peso (suddiviso in 4 intervalli).

peso in kg.	altezza in centimetri				
	fino a 150	tra 151/160	tra 161/170	tra 171/180	più di 180
fino a 60	$A_{11}=6$	$A_{12}=2$	$A_{13}=1$	$A_{14}=0$	$A_{15}=0$
tra 61 e 70	$A_{21}=3$	$A_{22}=9$	$A_{23}=14$	$A_{24}=7$	$A_{25}=0$
tra 71 e 80	$A_{31}=1$	$A_{32}=4$	$A_{33}=16$	$A_{34}=25$	$A_{35}=8$
più di 80	$A_{41}=0$	$A_{42}=0$	$A_{43}=4$	$A_{44}=6$	$A_{45}=15$

Fig. 6 - Esempio di matrice a due dimensioni.

Caratteristica essenziale della matrice è che l'accesso ad un suo elemento qualsiasi avviene tramite la n -pla di indici associata a questo (la coppia di indici nell'esempio di figura 6) senza passare attraverso operazioni di ricerca sequenziale. Questa caratteristica fa sì che una matrice ad una dimensione (vettore) si distingua da una lista lineare.

5. Tavole

Con il nome di « tavola » o di « tabella » si intende un insieme finito di elementi, ciascuno dei quali costituito da una coppia ordinata di dati, il primo dei quali viene chiamato « nome » o « chiave » dell'elemento, mentre il secondo è chiamato « valore » ed è costituito da informazioni associate alla chiave. Ad un elemento di una tavola si vuole accedere tramite la chiave.

Le tavole si presentano con frequenza elevatissima nei problemi di elaborazione automatica dei dati. Un esempio molto ovvio è costituito da un listino prezzi in cui il codice articolo costituisce la chiave, mentre il valore può comprendere il prezzo unitario dell'articolo e una descrizione di questo, come è illustrato nella figura 7.

Un altro esempio molto semplice può essere fornito da una tavola in cui ogni elemento è costituito dai vocaboli che individuano lo stesso oggetto in due lingue diverse, come nella fig. 8. In tale esempio il nome italiano ricopre il ruolo di chiave, se la tavola è utilizzata per una traduzione dall'italiano in inglese, mentre il ruolo di chiave è ricoperto dal nome inglese in caso di traduzione dall'inglese in italiano. Ciascuna delle due parti di cui si compone un vocabolario italiano-inglese e inglese-italiano è quindi una tavola in cui gli elementi sono ordinati secondo l'ordine alfabetico delle chiavi.

Naturalmente le tavole trovano numerosissime altre applicazioni, per esempio nella realizzazione dei compilatori e più in generale in tutti i casi di corrispondenze tra insieme non esprimibili, o esprimibili in modo poco conveniente, mediante relazioni matematiche.

(chiave)	(valore)	
Codice articolo	Prezzo unitario	Descrizione
61305 T	1200	Pialla regolabile MM 30
61320 R	1450	Saldatore elettrico
61340 M	2400	Cacciavite automatico
61356 C	2300	Seghetto alternativo
61357 A	14500	Trapano elettrico
61358 R	2700	Supporto per trapano
61386 T	4150	Levigatrice
61396 R	1950	Albero flessibile
61401 M	3600	Sega circolare
61410 V	6750	Sega traforo
61412 D	1400	Morsa da banco
61413 S	2100	Trapano a mano
62201 T	2400	Pialla regolabile MM 50
62202 V	1450	Mola MM 100
62204 D	2100	Mola MM 120
62207 M	4300	Saldatore a pistola
62262 S	13500	Pistola a spruzzo
62277 T	1600	Forbice per lamiera
62321 X	1100	Calibro in acciaio
62324 T	850	Pinza universale

Fig. 7 - Un esempio di tavola (listino prezzi).

6. Grafi e alberi

Il termine « grafo » (almeno secondo una definizione non rigorosa) viene impiegato per indicare una figura costituita da un insieme finito di punti, detti « nodi » o « vertici », e da segmenti o archi, detti « spigoli », che congiungono un certo numero di coppie di nodi.

Se ogni nodo viene considerato il supporto di un dato e gli spigoli sono considerati come la rappresentazione di una relazione tra i dati contenuti nei nodi da essi uniti, ed eventualmente a loro volta supporto di dati, allora il grafo può essere visto come la rappresentazione di una struttura astratta di dati, alla quale si assegna lo stesso nome.

Fig. 8 - Un esempio di tavola di traduzione.

Nome italiano	Nome inglese
Airone	Heron
Anitra	Duck
Beccaccia	Woodcock
Falco	Hawk
Gabbiano	Gull
Gufo	Owl
Oca	Goose
Passero	Sparrow
Pettiroso	Robin
Rondine	Swallow
Scricciolo	Wren
Tordo	Thrush
Usignolo	Nightingale

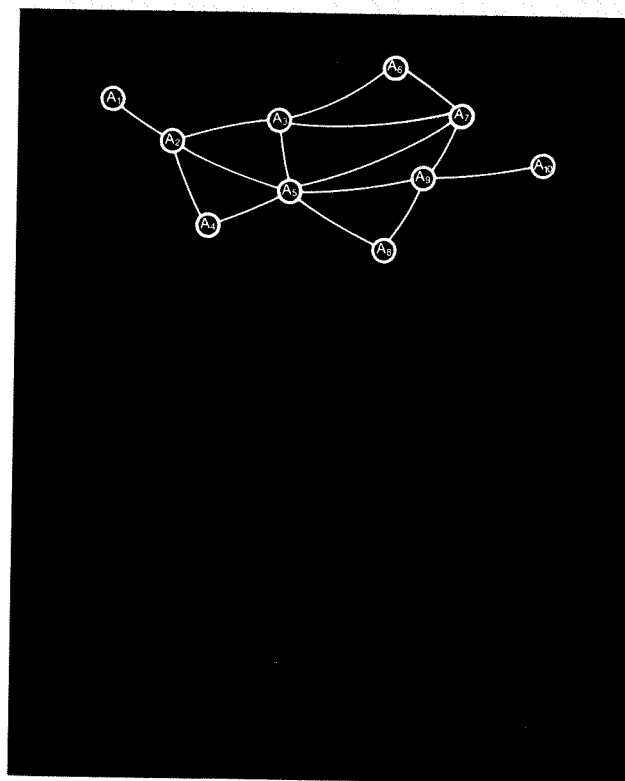


Fig. 9 - Un esempio di grafo.

La fig. 9 rappresenta ad esempio un grafo i cui nodi sono indicati con $A_1, A_2, \dots, A_{10}$; gli spigoli del grafo possono essere individuati mediante i nomi dei nodi da essi congiunti (così lo spigolo che congiunge il nodo A_3 con il nodo A_5 sarà chiamato spigolo $A_3 A_5$). Due nodi di un grafo si dicono « adiacenti » se esiste uno spigolo che li congiunge; una successione di nodi adiacenti prende il nome di « cammino » se non contiene due volte uno stesso spigolo.

In particolare si chiama « cammino semplice » un cammino comprendente nodi tutti distinti, ad eccezione eventualmente del primo e dell'ultimo che possono coincidere: in quest'ultimo caso il cammino semplice prende il nome di « ciclo ».

Se si fa riferimento alla figura 10, ad esempio, la successione $A_1, A_2, A_5, A_7, A_9, A_{10}$ costituisce un cammino semplice (fig. 10 a), il cammino semplice $A_5, A_3, A_7, A_9, A_8, A_5$ è un ciclo (fig. 10 b), mentre la successione di nodi $A_8, A_5, A_2, A_3, A_5, A_4$ è un cammino non semplice (fig. 10 c).

Infine se, presi due nodi arbitrari di un grafo, è sempre possibile congiungere tali nodi mediante un cammino, il grafo si dice « connesso ».

La nozione di grafo può dare origine alla nozione di « grafo orientato » in cui ogni spigolo è dotato di un orientamento che precisa la relazione logica che intercede tra i dati contenuti nei due nodi congiunti dallo spigolo stesso. Gli spigoli di un grafo orientato vengono rappresentati mediante segmenti o archi terminanti con una freccia, come è indicato nella fig. 11. Un particolare tipo di grafo connesso è costituito dall'« albero », caratterizzato dall'assenza di cicli. Per

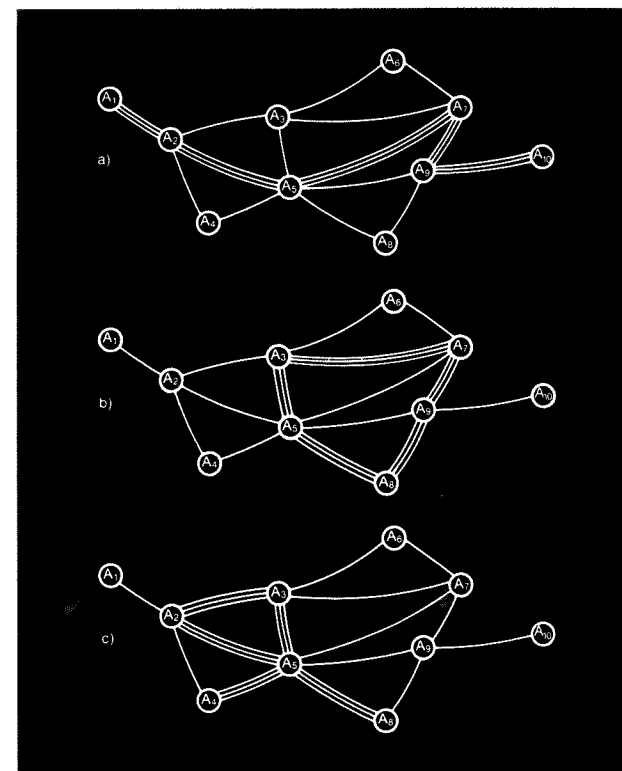


Fig. 10 - Alcuni esempi di cammini in un grafo.

illustrare più semplicemente le proprietà di un albero, consideriamo un albero orientato.

Si può verificare infatti che un albero orientato, fatta una semplice ipotesi sull'orientamento, possiede un nodo, detto « radice » dal quale non ha origine alcun arco orientato, mentre da ogni suo altro nodo ha inizio uno ed un solo arco orientato che è inizio di un cammino orientato conducente alla radice. I nodi dell'albero che sono origine, ma non estremi di arrivo di archi orientati prendono il nome di « foglie ». Ogni nodo può essere considerato radice del « sottoalbero » costituito dai nodi appartenenti a cammini orientati che terminano nel nodo considerato; le foglie possono in tal caso essere considerate come un caso limite di sottoalbero.

La fig. 12 rappresenta un particolare albero, detto « albero binario » nel quale ad ogni nodo, escluse ovviamente le foglie, giungono sistematicamente due archi orientati.

Grafi e alberi trovano applicazione nella risoluzione di problemi in cui i dati hanno una struttura logica appropriata.

Un grafo può ad esempio essere utilizzato per rappresentare un sistema di vie di comunicazione o in particolare un labirinto e quindi per risolvere il problema della determinazione di un cammino, eventualmente semplice, che congiunge due nodi.

Se poi ad ogni lato del grafo si associa un dato che esprime la lunghezza del lato stesso, cioè la distanza tra i due nodi, si può affrontare il problema della determinazione della distanza minima tra due nodi qualsiasi.

Un famoso problema che può essere affrontato con

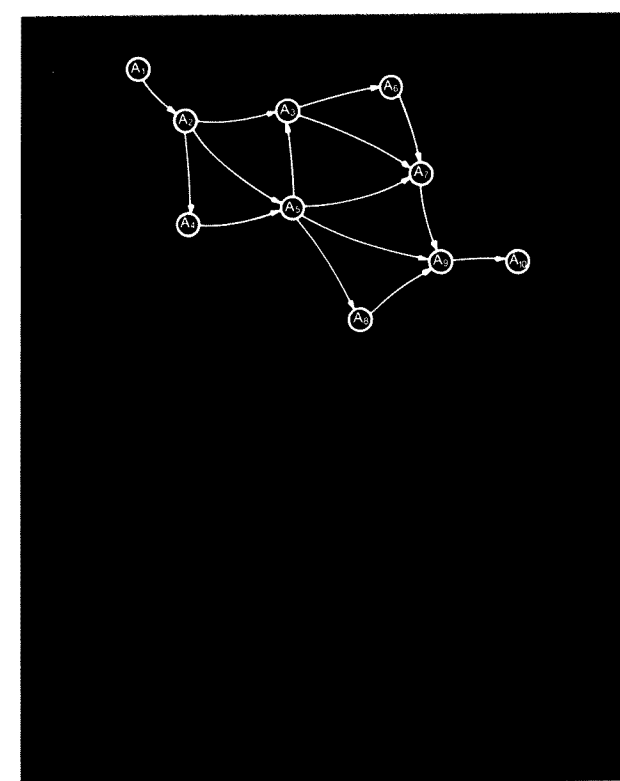


Fig. 11 - Un esempio di grafo orientato.

la rappresentazione dei dati mediante un grafo, è quello detto dei « ponti di Königsberg » che anzi è all'origine della teoria dei grafi.

Il problema, risolto da Eulero nel 1736, richiedeva di determinare se era possibile effettuare una passeggiata con ritorno al punto di partenza dopo aver attraversato una e una sola volta ciascuno dei sette ponti della città di Königsberg, costruita sulle rive e su due isole di un fiume secondo la pianta rappresentata schematicamente nella figura 13.

La soluzione del problema può essere ottenuta rappresentando le quattro zone della città come i nodi (1), (2), (3), (4) di un grafo e i sette ponti come spigoli che congiungono i quattro nodi nel modo indicato nella figura 14.

A questo punto basta chiedersi se esiste nel grafo così costruito un cammino chiuso cui appartengono una e una sola volta i sette spigoli rappresentanti i ponti. Ma un cammino chiuso con queste caratteristiche dovrebbe entrare in ogni nodo tante volte quante ne esce e quindi ad ogni nodo dovrebbe arrivare un numero pari di spigoli, proprietà non verificata dal grafo in esame.

Ovviamente Eulero non si limitò a risolvere il problema dei ponti di Königsberg, ma risolse un problema più generale, cioè quello della determinazione delle condizioni per l'esistenza in un grafo di un ciclo che percorra tutti gli spigoli una e una sola volta.

Correlato al problema di Eulero è il problema di Hamilton, che consiste nella determinazione di un ciclo che passa per tutti i nodi del grafo una e una sola volta.

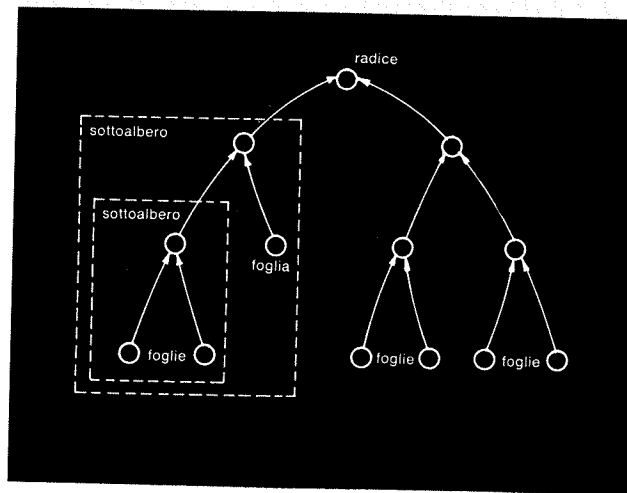


Fig. 12 - Un esempio di albero binario.

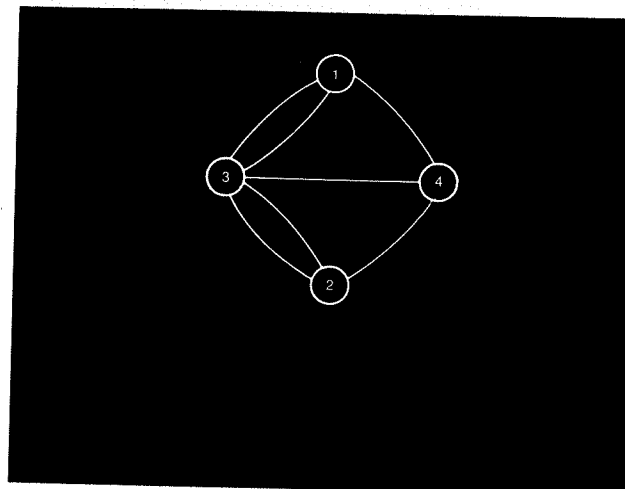
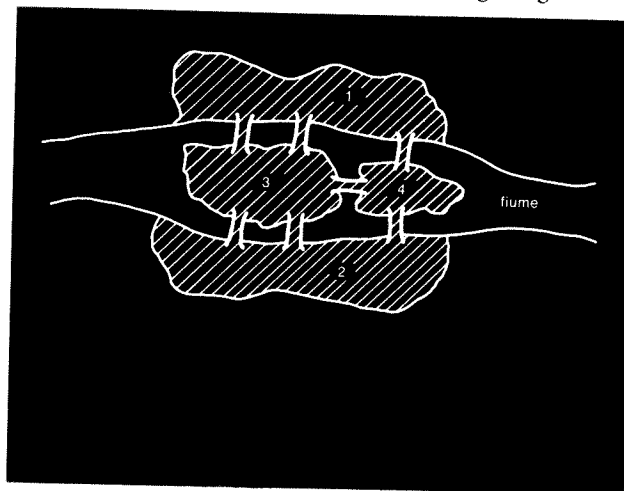


Fig. 14 - Grafo che rappresenta il problema dei sette ponti.

Qualche somiglianza con il problema di Hamilton ha poi il noto « problema del commesso viaggiatore » che prevede la ricerca del cammino più breve, o più economico, tra quelli che toccano un certo numero di località, rappresentabili come vertici di un grafo. Ovviamente i grafi possono essere utilizzati in moltissimi altri problemi, da quello classico dei colori nelle carte geografiche alla teoria dei giochi e allo studio delle reti elettriche. Un cenno particolare è opportuno a proposito dell'impiego dei grafi per la rappresentazione di un complesso sistema di attività, ciascuna delle quali può essere condizionata al completamento di altre.

Questa applicazione è divenuta molto frequente in questi ultimi anni e la tecnica impiegata per determinare gli elementi rilevanti del sistema di attività, quali il cosiddetto « cammino critico », che determina il tempo minimo intercorrente tra l'inizio della prima attività del sistema e il completamento dell'ultima, è usualmente nota sotto il nome di tecnica PERT (Program Evaluation and Review Technique). Per quanto riguarda gli alberi, se si trascurano le applicazioni del tutto ovvie, come la rappresentazione di un albero genealogico, essi possono essere im-

Fig. 13 - Il problema dei ponti di Königsberg.



piegati nella teoria dei giochi per rappresentare le possibili partite di un gioco di abilità e trovano una applicazione interessante nella rappresentazione delle strutture sintattiche delle frasi dei linguaggi di programmazione e in particolare delle espressioni aritmetiche.

Ad esempio, l'espressione aritmetica

$$A \cdot B + C \cdot (D \cdot E - F \cdot G)$$

può essere rappresentata convenientemente in forma grafica mediante l'albero binario della fig. 15 che mette in evidenza l'ordine in cui le operazioni devono essere eseguite e si presta bene per essere utilizzato dai compilatori.

A tale proposito è opportuno dare un cenno al problema dell'« attraversamento » di un albero, consistente nella determinazione di un algoritmo che permetta di percorrere i nodi dell'albero secondo un ordine che fa esaminare una e una sola volta ciascun nodo.

Se l'albero considerato è un albero binario, l'attraversamento può essere realizzato circondando l'albero con un percorso antiorario a partire dalla radice ed eseguendo in presenza di ciascun nodo le seguenti scelte:

- 1) se il nome è una foglia, viene esaminato;
- 2) se il nodo non è una foglia e ad esso si giunge dall'alto, viene ignorato;
- 3) se il nodo non è una foglia e ad esso si arriva dal basso per la prima volta, viene ignorato;
- 4) se il nodo non è una foglia e ad esso si arriva dal basso per la seconda volta, viene esaminato.

L'algoritmo ora descritto, che permette un attraversamento detto « post-ordinato », non è l'unico possibile. La sua applicazione all'albero della fig. 15 è illustrata nella figura 16; se si suppone che l'effetto dell'esame di un nodo sia la trascrizione del suo contenuto in una lista lineare, il risultato dell'attraversamento è in questo caso costituito dalla stringa:

$$AB \cdot CDE \cdot FG \cdot - \cdot +$$

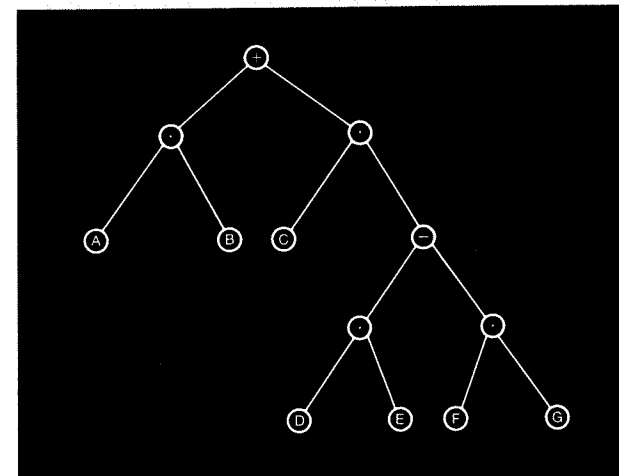


Fig. 15 - Rappresentazione di una espressione aritmetica mediante un albero binario.

che è la cosiddetta « rappresentazione polacca post-fissa » dell'espressione aritmetica già vista.

Tale rappresentazione, in merito alla quale sarebbe necessario un discorso a parte, è caratterizzata dal fatto che, scorrendo da sinistra verso destra l'espressione, ogniquale si incontra un operatore questo deve essere applicato ai due operandi che immediatamente lo precedono. La terna (due operatori ed operando) deve esser sostituita con il risultato dell'operazione e la scansione dell'espressione deve essere ripresa, sempre da sinistra verso destra, dal punto in cui era stata sospesa.

La rappresentazione polacca permette pertanto di ignorare la priorità associata ad ogni operatore aritmetico ed elimina l'uso delle parentesi.

7. Conclusione

La rassegna, molto succinta, che è stata svolta fin qui, ha permesso di evidenziare la strutturazione in insiemi dei dati che intervengono in una elaborazione e di esaminare i tipi più comuni di strutture.

Si pone ora il problema di rappresentare tali strutture informative nella memoria di un calcolatore, la quale è sostanzialmente costituita da una successione di celle, a ciascuna delle quali è associato un indi-

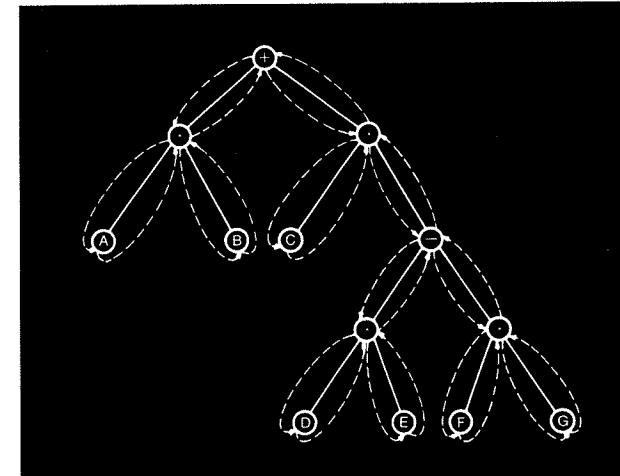


Fig. 16 - Attraversamento di un albero binario.

rizzo. Allo scopo di rendere agevole l'accesso dei dati appartenenti a strutture informative di una certa complessità memorizzate nella struttura elementare della memoria, è opportuno sovrapporre a questa una organizzazione più complessa.

Questa organizzazione, che viene realizzata per mezzo di programmi, viene detta « struttura concreta di memoria ». I più diffusi tipi di strutture concrete di memoria saranno oggetto di una rassegna in un successivo articolo che comparirà in questa stessa rivista.

8. Bibliografia

- [1] M.C. HARRISON, *Data Structures and Programming*, Courant Institute of Mathematical Sciences, New York University, 1972.
- [2] F.R.A. HOPGOOD, *Compiling techniques*, McDonald-Elsevier Computer Monograph, 1967.
- [3] M. ITALIANI, G. SERAZZI, *Elementi di Informatica*, Etas Libri, 1974.
- [4] D.E. KNUTH, *The art of computer programming*, Addison Wesley, 1969.
- [5] F. LUCCIO, *Strutture Linguaggi Sintassi*, Boringhieri Editore, 1974.
- [6] L. MURACCHINI, *Introduzione alla teoria dei grafi*, Boringhieri Editore, 1967.
- [7] O. ORE, *I grafi e le loro applicazioni*, Zanichelli, 1972.

Problemi nella interazione uomo-calcolatore

K.D. EASON, L. DAMORADAN, T.F. STEWART

University of Technology, Loughborough (GB)
Department of Ergonomics and Cybernetics

1. Introduzione

L'interfaccia tra l'uomo ed il calcolatore costituisce un punto cruciale per l'utilizzazione di questo mezzo, soprattutto se l'utente non è un esperto di calcolatori, ma usa la macchina come uno strumento di lavoro.

Gli autori hanno studiato questo problema attraverso una indagine condotta, in Inghilterra, su un ampio campione di utenti (254 utenti di 26 sistemi di elaborazione), appartenenti a diverse categorie professionali: impiegati, managers, specialisti.

In questo articolo vengono esposti i criteri e le conclusioni dell'indagine, che è stata condotta col concorso del Social Science Research Council inglese.

2. Ricerche sull'interfaccia

Gli specialisti di studi sul fattore umano cominciarono a prendere seriamente in considerazione l'interfaccia fra uomo e calcolatore, quando tale interazione poté aver luogo per mezzo di terminali « on line ». La presenza di un elemento di « hardware » voleva dire che l'interazione uomo-calcolatore poteva essere considerata come un caso particolare dell'interazione uomo-macchina. Era perciò evidente che per questo nuovo campo di studi occorrevo ricerche preliminari sulle immagini visuali che la macchina fornisce all'utente, e sui comandi con cui l'utente aziona la macchina. In seguito a tale interessamento, molto lavoro è stato compiuto sui dispositivi destinati all'uscita nell'interazione uomo-calcolatore (stampanti, visualizzatori alfanumerici, tracciatori di grafici, ecc.) e sui dispositivi destinati all'ingresso, come tastiere speciali, penne a luce, e comandi a « cloche ».

Una volta posto il problema di ottenere una efficace interazione fra uomo e calcolatore, si vide ben presto che ciò presenta difficoltà che non si trovano nell'interazione uomo-macchina. Nickerson (1969), per esempio, ha fatto notare le molte difficoltà relative al fattore umano associate al linguaggio del calcolatore,

e Miller (1968), ha sottolineato l'importanza di certe variabili del sistema, come il tempo di risposta del calcolatore.

La letteratura sui problemi umani associati con l'interfaccia fra uomo e calcolatore sta rapidamente crescendo ma è molto frammentaria. Un autore si interessa degli errori nell'introduzione dei dati, un altro della disposizione delle tastiere, un altro ancora della visualizzazione dei caratteri. Ne risulta qualcosa che assomiglia più a un catalogo che a un insieme integrato di conoscenze, e ancora meno a una teoria dell'interazione uomo-calcolatore.

Una delle ragioni di un tale stato di cose è che l'interazione uomo-calcolatore può assumere molte forme, e i relativi problemi non sono gli stessi per le diverse forme. Il calcolatore è un elaboratore di informazioni di uso generale, e il suo impiego può andare dalla registrazione di una singola informazione fino alla soluzione del più complesso problema matematico. Sotto molti aspetti, gli utenti che interagiscono col calcolatore su tutta questa gamma di mansioni, è come se interagissero con macchine differenti. La variabilità delle interazioni fra uomo e calcolatore è paragonabile a quella che passa fra l'andare in bicicletta e pilotare un aereo di linea. I problemi che riguardano l'utente in una situazione non sono quelli che si pongono in un'altra situazione.

Noi crediamo che, per giungere a rendere coerenti le conoscenze attuali sui problemi di interfaccia, occorre stabilire una classificazione delle diverse forme di interazione uomo-calcolatore. Questo scritto propone alcune classificazioni che abbiamo sviluppato mentre studiavamo le forme di interazione usate dagli utenti di 26 sistemi di elaborazione in Gran Bretagna.

Non abbiamo ritenuto possibile includere in un solo studio tutte le forme di interazione, e perciò abbiamo limitato l'indagine alle forme usate da una estesa categoria di utenti. Questi sono gli *utenti ignari del calcolatore*, che noi definiamo come individui operanti in una organizzazione, che non sono esperti nel-

la tecnologia del calcolatore, ma lo usano come strumento ausiliario nel loro lavoro. Perciò non si considera qui l'interazione fra programmatori e calcolatori. Gli obiettivi di questa inchiesta sono di accertare gli scopi per cui gli utenti ignari interagiscono con il calcolatore, le forme di interazione impiegate, e i problemi associati con ciascuna forma di interazione.

3. Classificazione delle forme di interazione

Se uomo e calcolatore debbono interagire, alcune condizioni debbono essere soddisfatte. Anzitutto deve essere possibile all'uomo di comunicare con il calcolatore usando i dispositivi di ingresso (vista, udito) e i dispositivi di uscita (mani, voce) che sono a sua disposizione. In secondo luogo, perché il significato della comunicazione sia compreso ad ambedue le estremità del canale, l'uomo e il calcolatore debbono avere in comune un sistema simbolico di codificazione. Infine, se si ammette che si possa realizzare un sistema di calcolo che adempia a tale condizione, occorre anche assicurarsi che l'uomo sia in grado di acquistare le necessarie abilità e di imparare le procedure atte a fare funzionare il sistema.

Questo modo di descrivere l'interazione fra uomo e calcolatore, che pone l'accento più sul punto di vista dell'utente, che su quello del progettista del sistema, ci porta ad identificare le tre componenti della interazione indicate qui di seguito, che sono state usate per classificare le forme di interazione incontrate nell'inchiesta.

a) Il mezzo di interazione

Abbiamo chiamato « mezzo di interazione » l'insieme delle caratteristiche fisiche della comunicazione fra uomo e calcolatore. Nell'inchiesta abbiamo fatto una distinzione di base fra utenti che hanno a disposizione un terminale « on line » di un calcolatore, e quelli che usano mezzi di comunicazione meno diretti. I terminali sono poi stati suddivisi secondo il tipo di uscita (copia a stampa, o visualizzazione) e il tipo di dispositivo di ingresso usato (tipo di tastiera, penna a luce, ecc.). L'interazione senza terminali consisteva tipicamente in fogli di codificazione o moduli di presentazione dati, per l'ingresso, e in copie a stampa per l'uscita. Tuttavia alcuni utenti non avevano alcun contatto con nessuno di tali mezzi, perché usavano altre persone come « interfaccia » col calcolatore, per preparare i dati per l'ingresso e/o interpretare le uscite del sistema.

b) Il modo di interazione

Abbiamo chiamato « modo » d'interazione la natura della comunicazione simbolica fra uomo e calcolatore. La caratteristica specifica dei modi di interazione usati dagli utenti ignari è quella di non permettere che un limitato accesso alla potenzialità del calcolatore. La maggior parte degli utenti impiegava un insieme di procedure appositamente studiate per permettere loro di effettuare interazioni relative al loro lavoro, ma che non davano accesso alle altre pos-

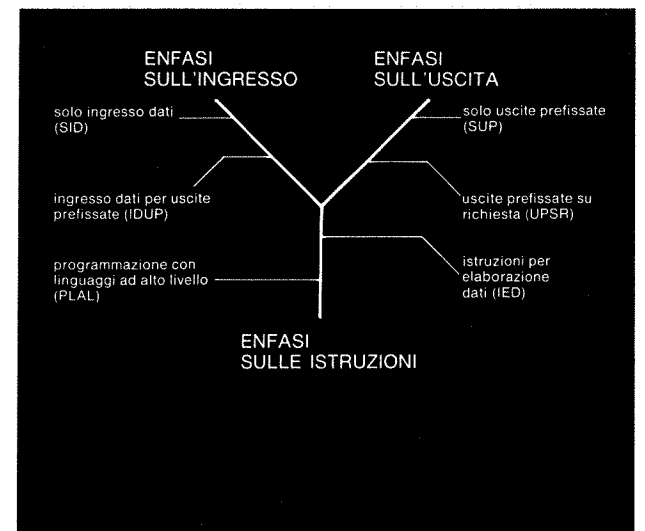


Fig. 1 - Modi di interazione.

sibilità che avrebbero potuto essere fornite dalla tecnologia del calcolatore. Alcuni di tali modi erano più restrittivi di altri, e noi abbiamo stabilito la classificazione di fig. 1 per rappresentare le differenze fra ampiezza e flessibilità di tali differenti modi.

In fig. 1 i modi di interazione sono divisi in tre gruppi: quelli che danno importanza soprattutto all'Ingresso dei dati, quelli che danno importanza alla ricezione della Uscita, e quelli che si occupano di istruire il calcolatore sui metodi secondo cui l'informazione va elaborata. In ogni gruppo sono stati rilevati due modi. Nel gruppo « Ingresso » v'era un modo che limitava l'utente al solo ingresso dei dati (SID) ed un altro in cui l'utente introduceva i dati ma otteneva anche una uscita dal calcolatore (IDUP). Il gruppo « Uscita » aveva un modo in cui l'utente riceveva automaticamente le uscite (SUP), e un altro in cui l'utente poteva richiedere le uscite che lo interessavano (UPSR). Solo nel gruppo « Istruzioni » l'utente poteva effettivamente manipolare e cambiare i dati. Uno dei modi di questo gruppo (IED) forniva all'utente un sottoinsieme di processi che egli poteva applicare ai dati. Il limitato campo di possibilità di questo modo era notevolmente ampliato nel secondo modo (PLAL) in cui gli utenti usavano linguaggi ad alto livello, come FORTRAN o ALGOL. Si sarebbero potuti identificare modi di interazione di potenza ancora maggiore, come i linguaggi di macchina, ma chiunque sia capace di interazioni di questo tipo non può essere considerato un utente « ignaro ».

c) Assistenza all'utente

L'ultima componente della interazione è il metodo con cui l'utente acquista la conoscenza e l'abilità necessarie per comprendere ed usare utilmente il calcolatore. Noi abbiamo chiamato ciò « assistenza all'utente ».

Abbiamo analizzato in due modi le specie di assistenza di cui un utente ha bisogno. In primo luogo abbiamo fatto una analisi delle funzioni per le quali un utente abbisognava di assistenza. In certi casi, l'utente doveva imparare ad adoperare un termi-

nale, ed inoltre impadronirsi delle procedure per fare funzionare il sistema; oppure doveva poter giudicare quali delle possibilità del sistema erano adatte alla sua mansione, e doveva anche sapere cosa fare in caso di cattivo funzionamento, e come lasciare il sistema se si trovava in difficoltà. Doveva anche sapere a chi rivolgersi per aiuto e infine, se trovava che il sistema non poteva adempiere alle funzioni richieste, doveva essere capace di sottoporre le sue esigenze ai responsabili del funzionamento e dello sviluppo del sistema.

Il secondo modo di analisi riguardava la maniera con cui erano fornite tali forme di assistenza. In certi casi lo stesso sistema forniva una parte dell'assistenza necessaria. Ciò era ottenuto in due modi: provvedendo una documentazione di assistenza, o includendo l'assistenza nella interazione fra utente e sistema. Le funzioni più frequentemente servite in questi modi erano quelle che specificavano le procedure per usare il sistema, e le possibilità disponibili. Manuali di utente contenenti queste informazioni erano di uso comune: qualche sistema cercava di includere queste informazioni nel sistema stesso. Tuttavia nessuno dei sistemi investigati poteva soddisfare a tutte le esigenze di assistenza dell'utente; in tutti una parte dell'assistenza era fornita da membri del personale. Nell'esaminare i sistemi « umani » di assistenza abbiamo fatto una distinzione fondamentale fra meccanismi di assistenza formali ed informali. Alcuni dei sistemi esaminati avevano formalizzato dei servizi di assistenza, o dei gruppi di collegamento con gli utenti, che svolgevano molte di tali funzioni. Nella maggior parte dei sistemi, tuttavia, solo una o due funzioni erano formalizzate: generalmente l'addestramento (spesso solo all'atto della prima installazione del sistema), e la manutenzione e riparazione dell'« hardware ». Per altre funzioni, se pure si era provveduto ad esse, ciò era stato fatto con la creazione di una rete informale di contatti fra personale del sistema, utenti esperti, ed utenti nuovi e non esperti. In molti sistemi abbiamo trovato « esperti locali »: utenti che, in virtù di una lunga esperienza e di uno speciale interesse nel sistema, erano considerati dai loro colleghi come fonte di ogni conoscenza circa il sistema.

4. Tipi di Utente

Durante l'inchiesta sono stati intervistati in totale 254 utenti di calcolatori e, oltre a classificare le forme di interazione impiegate, abbiamo classificato gli utenti secondo il posto di lavoro occupato. Abbiamo diviso gli utenti in tre grandi categorie:

I) Utenti impiegati (n. 103). Questa categoria comprende una grande varietà di posti di lavoro, come, per esempio, fatturatori, contabili, magazzinieri e cassieri di banca. Tutti questi impieghi erano caratterizzati dal loro carattere di « routine » ripetitiva, e da un uso del calcolatore relativamente ben definito.

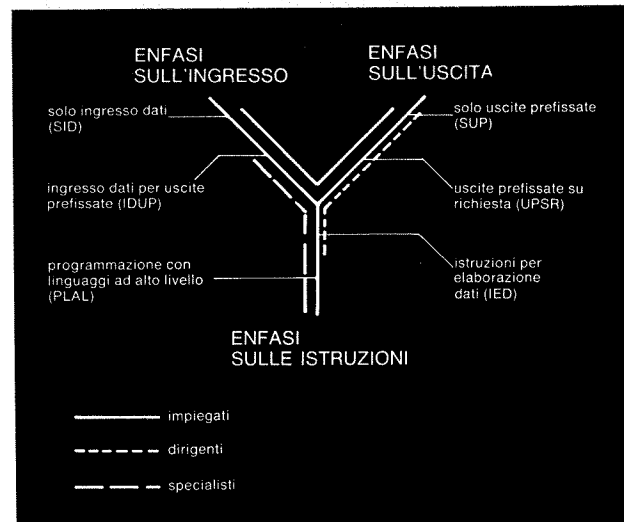


Fig. 2 - Modi di interazione.

II) Utenti dirigenti (n. 82). I dirigenti utenti del calcolatore andavano dai direttori generali, attraverso livelli intermedi di dirigenza, fino alla dirigenza e supervisione di prima linea. Il campione esaminato comprende soprattutto dirigenti di prima linea e di livello intermedio, in parte per l'inaccessibilità degli alti dirigenti, in parte perché molti di essi dichiararono di essere relativamente poco interessati dal calcolatore. Brady (1967), trovò che l'alta dirigenza negli U.S.A., similmente non era direttamente interessata dai sistemi di calcolo.

III) Utenti specialisti (n. 69). Questa categoria di utenti include le varie specie di addetti alla soluzione dei problemi che si presentano nell'industria e nel commercio; persone che potevano essere consultate su una disciplina specializzata, come ingegneria, progettazione, ricerca operativa, economia, statistica, ecc. Essi non erano specialisti di calcolatori, ma, per la loro specializzazione, avevano specifiche esigenze di aiuto dal calcolatore.

Si è trovata una forte correlazione fra modi di interazione e diversi tipi di utente, come indica la fig. 2. La maggior parte degli utenti impiegati usava modi di interazione orientati verso l'Ingresso, benché una importante minoranza usasse modi di Uscita.

A contrario i dirigenti erano interessati soprattutto ai modi di Uscita, benché ve ne fossero 24 che facevano uso di modelli valutativi, e quindi operassero secondo un limitato modo d'Istruzione. Gli specialisti predominavano nei modi di Istruzione, benché alcuni si limitassero alle versioni più interattive dei modi di Ingresso.

5. Problemi degli utenti con le forme di interazione

Nell'inchiesta, le interviste agli utenti consistevano di due parti: una analisi della forma di interazione usata dall'utente, e l'esame di una serie di argomenti che potevano influire sulla sua reazione al calcolatore. Tali argomenti erano delle seguenti quattro specie:

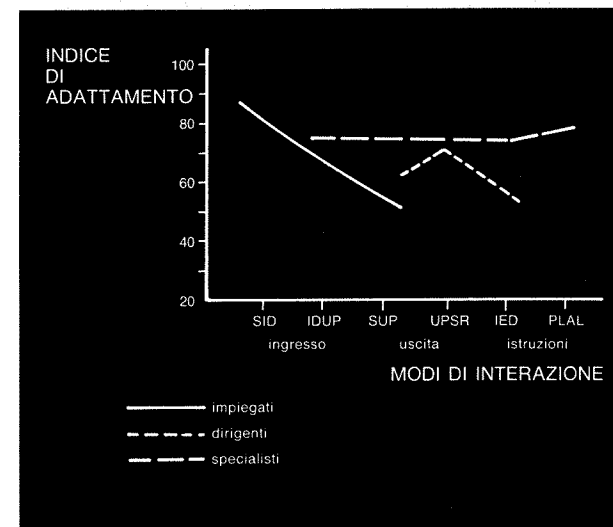


Fig. 3 - Effetto del modo di interazione sull'adattamento dell'utente.

a) Argomenti relativi all'adattamento alla mansione. Si esaminava quanto soddisfacentemente il sistema di calcolo provvedesse alle esigenze della mansione dell'utente. Questa sezione dell'intervista comprendeva domande circa la rilevanza, l'esattezza, la completezza e la tempestività delle informazioni fornite dal sistema.

b) Argomenti relativi alla facilità di uso. Questa sezione esaminava se le uscite del calcolatore potevano essere lette ed analizzate facilmente, e con quanta facilità si potesse fare funzionare il sistema. Se veniva usato un terminale, domande addizionali riguardavano l'uso del terminale, la disposizione del posto di lavoro, ecc.

c) Argomenti relativi all'assistenza all'utente. Questa sezione comprendeva due specie di domande: quelle riguardanti l'assistenza fornita all'utente durante il funzionamento del sistema, per esempio con consigli o in caso di guasto; e quelle riguardanti il ruolo dell'utente nel progetto e nello sviluppo del sistema.

d) Conseguenze indirette del sistema. Una grande varietà di domande esplorava i possibili diversi effetti dell'uso del calcolatore nei riguardi dell'utente, del suo lavoro, dell'organizzazione. Esse comprendevano, per esempio, la soddisfazione sul lavoro, le possibilità di avanzamento, lo sviluppo di nuove capacità, ecc.

L'inchiesta ha rivelato che mentre certi argomenti erano di importanza critica per certi utenti, gli stessi erano relativamente poco importanti per altri. I risultati sono presentati particolareggiatamente altrove (Eason et al., 1974): qui di seguito viene dato un breve resoconto dei risultati riguardo all'interfaccia uomo-calcolatore.

a) Il mezzo di interazione

120 utenti interagivano con il calcolatore mediante terminali « on-line ». Di questi utenti, 30 erano impiegati, 50 specialisti e solo 14 dirigenti. Ciò indica che, benché influenti autori, come Licklider, (1965),

e Morton (1968) abbiamo sottolineato le potenzialità dell'uso dei terminali da parte dei dirigenti, i terminali stessi non sono ancora stati generalmente accettati nell'ambiente dirigenziale. Nei casi in cui i dirigenti investigati avevano sperimentato l'uso dei terminali, s'era poi rivelata una forte tendenza ad abbandonare l'interazione diretta e a utilizzare una « interfaccia umana », che usasse il calcolatore per conto del dirigente.

Benché gli studi sul fattore umano nell'interazione uomo-calcolatore siano concentrati sul mezzo di interazione, pochi utenti hanno avuto grossi problemi nell'uso dei terminali. La difficoltà più spesso citata era quella della tastiera alfanumerica completa che, normalmente, con il suo misterioso schieramento di tasti speciali, appariva, al non dattilografo, uno strumento il cui maneggio incuteva un certo timore. Un fattore di scontento per quasi tutti gli utenti di terminali era la quasi universale mancanza di considerazione per il mobile associato al terminale. Abbiamo trovato ben pochi esempi di tentativi di disegnare un posto di lavoro completo per l'utente, con adeguato sedile, spazio di tavolo, alloggiamenti per manuali e documenti di lavoro.

b) I modi di interazione

Alcune delle più interessanti differenze fra utenti si riferivano ai modi di interazione adottati. I modi di Ingresso erano quelli che provocavano i commenti più sfavorevoli, per le loro conseguenze indirette. Questi modi erano usati soprattutto dagli impiegati e normalmente consistevano in cicli brevi ed operazioni ripetitive. Gli utenti si lamentavano di eccessiva « routine », di poca soddisfazione sul lavoro, e di essere comandati dal calcolatore. Questi problemi erano molto gravi per quegli utenti la cui interazione era limitata all'introduzione dei dati: gli stessi problemi, benché sempre presenti, erano meno sentiti da quegli utenti non dovevano solo introdurre i dati, ma che ricevevano inoltre dal calcolatore un servizio che permetteva loro di espletare altri compiti: così, per esempio, l'impiegato di banca che controlla il saldo di un conto corrente per conto di un cliente. C'erano due importanti differenze fra questa specie di lavoro e quello della semplice introduzione dei dati: in primo luogo il lavoro era più complesso e più variato, e in secondo luogo l'impiegato poteva considerare il calcolatore come il suo servo e non come il suo padrone.

I problemi degli utenti dei modi di Uscita e di Istruzione non riguardavano tanto le conseguenze indirette del sistema, quanto la sua capacità di soddisfare ai bisogni della mansione e la sua facilità d'uso. Per esaminare la relazione generale fra l'adattamento alla mansione e il modo di interazione abbiamo costruito un sistema di punteggio complessivo ricavato dai risultati di tutte le domande che riflettevano un problema riguardante la mansione. Tali punteggi sono dati in fig. 3, e sono espressi in percentuale delle domande le cui risposte non rivelano problemi significativi, rispetto al numero totale delle domande

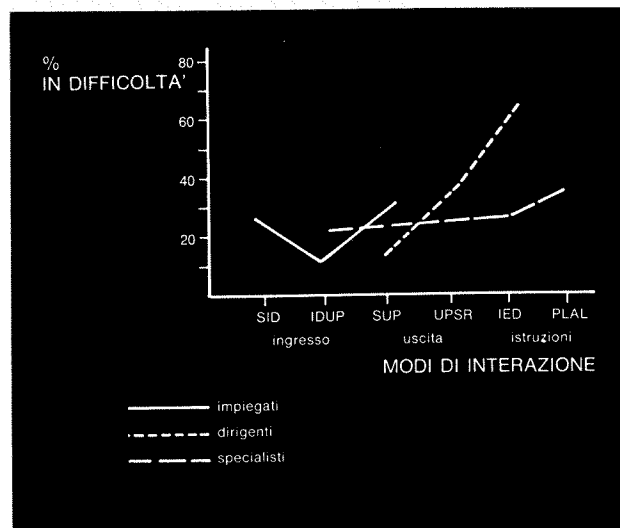


Fig. 4 - Reazioni degli utenti alle procedure operative di differenti modi di interazione. In ordinate, la percentuale di utenti che hanno avuto difficoltà con le procedure operative.

poste: perciò un punteggio alto rappresenta un buon grado di adattamento alla mansione.

Secondo la nostra ipotesi, l'adattamento alla mansione sarebbe dovuto migliorare con la maggiore potenza e flessibilità del modo di interazione: perché, quanta maggior potenza l'utente fosse in grado di comandare, tanto più facilmente avrebbe ottenuto i servizi che desiderava. Per confermare questa ipotesi, la fig. 3 avrebbe dovuto indicare un aumento dell'indice di adattamento da sinistra a destra: e questo non è il risultato ottenuto. Un'analisi dei risultati per ogni tipo di utente può spiegare la situazione. L'adattamento alla mansione per utenti impiegati diminuisce dai modi che riguardano l'ingresso a quelli che riguardano l'uscita. Ciò non sorprende perché il sistema è stato predisposto per facilitare appunto quella sola mansione che l'impiegato addetto all'ingresso dati deve svolgere. Nei modi IDUP e SUP l'impiegato era impegnato in mansioni che usavano l'uscita del calcolatore e l'adattamento alla mansione cadeva rapidamente quando venivano svolte mansioni differenti. Il risultato era quindi che spesso gli impiegati meno soddisfatti del servizio ricevuto erano proprio quelli addetti ai lavori che erano di maggiore soddisfazione, in quanto usavano l'uscita. L'adattamento alla mansione degli specialisti ha mostrato piccole variazioni qualunque fosse il modo usato. Ciò può sorprendere perché, per molti specialisti, la ragione per imparare un programma era quella di ottenere un migliore servizio. Benché il modo PLAL abbia mostrato un punteggio leggermente migliore degli altri modi, restava sempre uno scarico sostanziale fra gli adempimenti di cui gli specialisti avevano bisogno, e quelli realmente ottenuti.

Un quadro ancora più impressionante venne fornito dai risultati per i dirigenti. Nel modo SUP l'adattamento alla mansione è risultato basso tanto per i dirigenti che per gli impiegati. E' noto da qualche

tempo (p. es. vedi Ackoff 1968) che i sistemi di informazione per dirigenti, che operano tipicamente secondo tale modo, spesso forniscono informazioni irrilevanti e non aggiornate. L'utente ha ben poca influenza diretta sull'informazione che riceve e su quando la riceve. Nel modo UPSR l'utente ottiene una certa influenza su tali fattori. I risultati dell'inchiesta mostrarono, come previsto, che il modo UPSR dava agli utenti un miglior adattamento alla mansione che il modo SUP; tuttavia, il modo successivo, IED, che offriva maggior potenza e flessibilità, anziché continuare tale tendenza, portava a una diminuzione dell'adattamento.

Nonostante le manifeste differenze fra i tre tipi di utente, essi, nelle loro reazioni ai modi di interazione, hanno una caratteristica comune. I modi, che sono stati tipicamente previsti per il loro lavoro, non soddisfano le loro necessità: che sono necessità personali per gli impiegati, necessità di mansione nel caso di specialisti e dirigenti. Sfortunatamente, quando venivano procurati dei modi che potenzialmente avrebbero dovuto soddisfare queste necessità, difficilmente si sono ottenuti i risultati previsti.

Le ragioni di tale situazione sono molte, ma due sono di primaria importanza. In primo luogo, quanto più il sistema diventa complesso, tanto più è difficile da capire: e tuttavia, perché un dirigente o uno specialista possa usarlo efficacemente, bisogna che comprenda e accetti le premesse secondo le quali il sistema opera. Troppo spesso, come sottolinea Edstrom (1971) i sistemi sono progettati su basi normative, per ottenere ciò che il progettista ritiene necessario, invece di ciò di cui ha bisogno l'utente. Perciò l'utente, anche se può disporre di alcune scelte, non otterrà mai quello che cerca. Il secondo argomento riguarda la facilità di operazione dei sistemi flessibili. Come ha indicato Testa, (1973), spesso gli utenti trovano difficoltà a usare le istruzioni di impiego dei sistemi di calcolo. La fig. 4 dà le percentuali degli utenti intervistati che dichiararono di aver avuto difficoltà nel comprendere e nell'impiegare le procedure « software ».

Le procedure di « software » che gli utenti adoperavano avevano diverse funzioni: per introdurre dati, per fare ricerche nelle basi di dati, per specificare informazioni da estrarre da memoria, per specificare operazioni da eseguire su dati, ecc. Una analisi secondo i modi ha mostrato che gli utenti impiegati non incontravano grandi difficoltà, probabilmente perché essi usavano regolarmente il sistema e imparavano a fondo le procedure. Ciò era meno vero per gli utenti impiegati che usavano il modo SUP: essi erano, per la maggior parte, personale di grado superiore che faceva un uso meno frequente del sistema di calcolo, e, da questo punto di vista, si avvicinavano agli utenti dirigenti e specialisti che, essendo utenti solo occasionali, tendevano a non impadronirsi a fondo delle procedure che dovevano usare. Per gli specialisti si è notato un continuo aumento delle percentuali di difficoltà incontrate man mano che i

modi diventano più flessibili: per i dirigenti l'aumento era impressionante.

Questo è un aspetto particolare di un risultato generale rivelato dall'inchiesta. Gli utenti che cercavano un miglior servizio o un migliore impiego con l'usare un modo che fosse potenzialmente adatto alle loro necessità, si trovavano a scambiare un certo numero di problemi contro certi altri: e cioè con i problemi derivanti dall'uso di una tecnologia complessa e non familiare. Un interessante e importante risultato ausiliario è che, per ogni tipo di utente, questo problema diventava critico per modi differenti. Gli utenti specialisti, per esempio, trovavano soddisfacente per il loro lavoro il modo IED, mentre gli utenti dirigenti incontravano in questo molte difficoltà. Ne concludiamo che, in differenti tipi di utente, il rapporto operante fra sforzo mentale e vantaggi ottenuti è diverso: e cioè gli specialisti sono molto più disposti dei loro colleghi dirigenti a sostenere lo sforzo necessario per usare un sistema di calcolo. Il rapporto sforzo mentale-vantaggio in generale risulta essere stato spiegato nel modo migliore dal principio di Zipf (1965) secondo cui le persone adottano quella condotta che corrisponde al minimo tasso medio di lavoro probabile.

c) Assistenza all'utente

Nella maggior parte dei sistemi di calcolo, il sistema di assistenza all'utente in opera era un misto di formale e di informale, difficile da definire. Nei sistemi usati dagli impiegati la tendenza era di avere squadre di assistenza formali all'atto dell'installazione, che venivano sciolte in seguito. Però gli impiegati, più degli altri gruppi, esprimevano la necessità di assistenza: molto del loro lavoro era concentrato sul calcolatore, ed essi desideravano capire come lavorava, quale apporto essi gli davano, e in che relazione questo stava rispetto al funzionamento dell'intero sistema. Dispositivi formali che rispondessero a queste esigenze erano molto rari. L'assistenza richiesta da specialisti e dirigenti era più orientata verso la mansione, ed era richiesta per conoscere le possibilità disponibili e come usarle. Qualcuno dei sistemi per specialisti cominciava a riconoscere queste necessità e a provvedere squadre di « collegamento » e di « assistenza » all'utente: ma raramente ciò avveniva nei sistemi per dirigenti. Se non v'erano disposizioni formali per venire incontro a queste necessità, gli utenti tendevano ad affidarsi agli « esperti locali » per la guida richiesta. L'importanza di una rete locale di assistenza all'utente aumentava quando gli utenti pensavano di non potere, o non desideravano, occuparsi direttamente delle complicazioni del sistema e preferivano usare un intermediario.

6. Le conseguenze di una difettosa progettazione dell'interfaccia

Si può obiettare che i problemi, elencati con tanta varietà in questo studio, siano in fin dei conti di interesse puramente accademico, perché, in pratica,

Conseguenze	Tipo di utente	
	Impiegati	Dirigenti e Specialisti
Per il funzionam. del sistema	Adattamento dell'utente	1. Non uso 2. Limitato uso delle possibilità 3. Intermediario umano
Per l'utente	Frustrazione e abbassam. del morale	Incapacità di realizzare la potenzialità del sistema

Fig. 5 - Le conseguenze dei problemi di interfaccia per differenti tipi di utente.

gli utenti utilizzavano i sistemi. Vi sono però certi aspetti della reazione dell'utente a tali sistemi, che fanno pensare che questi problemi possano avere serie conseguenze. Tali aspetti sono indicati nella fig. 5.

Le conseguenze erano diverse per utenti impiegati e per utenti saltuari, come gli specialisti e i dirigenti. Generalmente, gli utenti impiegati non avevano altra scelta che di accettare le insufficienze del sistema e porvi rimedio alla meglio. L'adattamento dell'utente è un fenomeno che si incontra in ogni specie di sistemi tecnologici, quando l'uomo, che è flessibile, deve lavorare con tecnologie rigide: e i sistemi di calcolo non fanno eccezione. Il costo del dover continuamente compensare le deficienze del sistema stava nell'aumento della frustrazione dell'impiegato, il che conduceva all'abbassamento del suo morale, e alla perdita di interesse nel lavoro.

Gli utenti occasionali del calcolatore, il dirigente e lo specialista tolleravano assai meno un sistema difettoso. Essi si attendevano uno strumento adatto ai loro bisogni, e non erano disposti a modificare il loro comportamento per adattarsi alle sue necessità. Le conseguenze di una interfaccia difettosa per tali utenti erano differenti, e le reazioni erano di tre tipi:

1) Non uso del sistema. Una possibilità era che l'utente, dopo una prova iniziale del sistema, decidesse di non adoperarlo. A differenza degli utenti impiegati, gli utenti dirigenti e specialisti avevano spesso tale possibilità e senza dubbio molti ne hanno approfittato. Se v'erano utenti di questa specie fra i soggetti della nostra inchiesta, noi non li abbiamo rilevati perché abbiamo esaminato solo utenti attivi.

2) Uso limitato del sistema. Un vantaggio dei sistemi flessibili e interattivi era quello di offrire scelte. Molti utenti di tali sistemi hanno ammesso che essi tendevano a non usare che poche delle possibilità disponibili e a trascurare le altre. Si può obiettare che le altre possibilità potevano non essere interessanti per l'utente: ma questa reazione era tanto dif-

fusa da far pensare che molte possibilità convenienti non siano state usate.

3) L'uso di un intermediario umano. Una reazione particolarmente caratteristica dell'utente dirigente era l'usare qualcuno come intermediario umano fra l'utente e il sistema. In molti sistemi esaminati si prendevano misure per poter usare questo modo di interazione: e questo metodo è sempre più comune anche altrove. Martin (1973), per esempio, descrive un sistema in cui « segretarie di dati » operano in un « centro di informazione » per fornire le informazioni richieste da alti dirigenti. Per il dirigente tale metodo ha naturalmente molti vantaggi, dato che lo solleva da molti problemi relativi alla facilità d'uso. Non è facile sapere però se un tale sistema, per quanto ben organizzato, possa avere lo stesso risultato di una collaborazione « simbiotica » effettiva fra l'uomo e il calcolatore.

7. Verso una più efficace interazione fra uomo e calcolatore

Questo scritto si propone più di identificare i problemi incontrati dall'utente ignaro che di tentare di risolverli. Come avviene spesso, l'inchiesta rivolge la sua attenzione ad argomenti sui quali è necessaria una ulteriore ricerca, in particolare sulla dinamica del rapporto fra sforzo e vantaggio dell'utente, che sembra abbia una funzione decisiva nel determinare la reazione dell'utente al sistema. Nondimeno i risultati dell'inchiesta suggeriscono tre specie di argomenti che il progettista di sistemi deve tenere presente, se vuole ridurre lo sforzo che l'utente deve sostenere per raggiungere i suoi scopi, e se vuole rendere sempre maggiore l'interazione uomo-calcolatore.

1) L'interfaccia. I risultati dell'inchiesta rafforzano l'esigenza di progettare con maggior cura l'interfaccia « hardware » fra uomo e calcolatore, ma indicano anche l'importanza delle questioni riguardanti l'interfaccia « software ». E' interessante notare, per esempio, che mentre vi è un certo grado di standardizzazione delle tastiere, non v'è nulla di simile nelle procedure « software » con cui si effettuano operazioni basilari, come ingresso dati, introduzione dell'utente, messaggi di errore, ecc. Ne risulta che quando l'utente passa da un sistema ad un altro deve imparare un nuovo « linguaggio ».

2) Assistenza all'utente. I più flessibili e perfezionati sistemi che ora vengono posti a disposizione degli utenti, inevitabilmente esigono molto da essi. Anche se il funzionamento del sistema è reso facile, l'utente deve pur sempre conoscere le diverse opzioni a sua disposizione, capire che cosa ognuna di esse può fare, apprezzare la loro validità per lo scopo che si propone, e sapere come usarle. Sarà ben raramente possibile addestrare l'utente ad apprezzare pienamente la potenzialità di un sistema, e spesso l'apprendimento avrà luogo durante l'uso. Questo processo richiede che l'utente sia sufficientemente assistito fornendogli documentazione e consulenti umani. Al

presente ciò avviene più per caso che a ragion veduta.

3) Interessamento dell'utente alla progettazione del sistema. La quantità di sforzo che un utente deve sostenere per usare un sistema dipende in parte dalla sua precedente esperienza di esso, e in parte dalla misura in cui esso fornisce una risposta valida alle sue esigenze, in una terminologia a lui familiare. Il miglior modo di assicurarsi che, quando il sistema è in servizio, tutte queste forme di sforzo sono minimizzate, è di interessare l'utente al progetto del sistema. L'interessamento dell'utente può assicurare che egli riceva il servizio di cui abbisogna, che egli comprenda il servizio che può ricevere, che egli sappia far funzionare il sistema, e che egli sia impegnato nel sistema prima che esso sia installato, e non lo guardi con sospetto e dubbio.

8. Conclusioni

I risultati dell'inchiesta mostrano che attualmente i progettisti di sistemi creano interfacce di utente che permettono solo un limitato accesso alla potenzialità del calcolatore. Con tali limitazioni l'interfaccia può essere resa di uso abbastanza facile per l'utente ignaro. Tuttavia, in molti casi, le forme di interazione più semplici sono insufficienti. Esse, o contribuiscono a creare occupazioni ripetitive, come nel caso di utenti impiegati, oppure non soddisfano a sufficienza le esigenze relative alla mansione dell'utente, perché non sono abbastanza flessibili e potenti.

Ogni gruppo di utente necessita perciò di più complesse forme di interazione. Sfortunatamente le forme più complesse sono spesso di difficile uso, e l'utente non può o non vuole fare lo sforzo necessario per impadronirsi delle procedure « software » del sistema. Ne risulta, o una restrizione autoimposta della gamma di possibilità utilizzate, o la necessità di affidarsi a esseri umani che agiscano come intermediari. Man mano che le forme di interazione diventano più complesse, aumenta la possibilità che il calcolatore possa fornire un aiuto importante ad ogni tipo di utente. Ma l'aumento di complessità conduce a problemi critici di interfaccia, che debbono essere risolti se si vuole realizzare l'intero potenziale del calcolatore.

9. Bibliografia

- [1] ACKOFF, RUSSELL L. (1968), *Management Misinformation Systems*, Management Decision 2:1, 4-8.
- [2] BRADY, RODNEY H. (1967), *Computers in Top-Level Decision Making*, Harvard Business Review, 45; 4: 67-76.
- [3] EASON, KENNETH D., DAMODARAN, LEELA, and STEWART, THOMAS, F.M. (1974), *A Survey of Man-Computer Interaction in Commercial Applications*, Forthcoming Social Science Research Council Report, London.
- [4] EDSTROM, ANDERS (1971), *Determining Management Information Requirements*, relazione al International Seminar on Organisations and Computers, Dubrovnik, Yugoslavia, 16-18 August.
- [5] LICKLIDER, J.C.R. (1965), *Man-Computer Partnership*, International Science and Technology, 41: 13-26.

[6] MARTIN, JAMES (1973), *Design of Man-Computer Dialogues*, New Jersey, Prentice Hall.

[7] MILLER, ROBERT B. (1968), *Response Time in Man-Computer Conversational Transactions*, Proceedings of the Fall Joint Computer Conference, 267-277.

[8] MORTON, MICHAEL S. (1968), *Visual Display Systems and Management Problem Solving*, European Business, April, 37-46.

[9] NICKERSON, R.S. (1969), *Man-Computer Interaction: A Challenge for Human Factors Research*, Ergonomics, 12: 4: 501-517.

[10] TESTA, CHARLES J. (1973), *The Evolution of Man-Computer Symbiosis*, Computers and Automation, 22: 5: 11-12.

[11] ZIPF, GEORGE K. (1965), *Human Behaviour and the Principle of Least Effort*, New York: Hafner, Second edition.

Metodi automatici di gestione delle scorte nelle aziende di distribuzione

ROBERTO BRESCHI

Honeywell Information Systems Italia
Servizio Applicazioni

1. Introduzione

Con l'estendersi e l'affinarsi della automazione aziendale si sono sviluppate tecniche volte alla gestione delle scorte, per le quali negli anni 60 ci si limitava solitamente ad una contabilità delle entrate e uscite. La tenuta del giornale di magazzino e l'aggiornamento più o meno frequente dei saldi contabili erano infatti elaborazioni classiche (come la fatturazione e le paghe e stipendi) nate con le prime applicazioni del calcolatore in azienda.

Oggi il problema si è spostato sui temi della gestione di un saldo permanente (sempre più spesso con utilizzo di terminali per introduzione dati ed interrogazione degli archivi) e del controllo automatico delle giacenze con emissione delle proposte di riapprovvigionamento.

In alcuni casi questi obiettivi sono già stati raggiunti, ma per la maggioranza delle aziende questo problema è ancora nella lista dei lavori che il servizio EDP si propone di automatizzare in un futuro più o meno prossimo.

La problematica della gestione delle scorte si pone in termini diversi per magazzini di materie prime o semilavorati che alimentano linee di produzione e per magazzini di articoli destinati alla distribuzione. In questa sede noi ci proponiamo di trattare il secondo tipo di magazzini.

E' chiaro come in ogni caso una automazione della gestione si presenta come una delle più allettanti per i vantaggi economici ricavabili.

Infatti da un lato scorte elevate significano:

- pesanti immobilizzi finanziari
- alte spese di immagazzinaggio (locali, uomini, servizi)
- possibilità di deterioramento fisico
- possibilità di obsolescenza tecnica
- rischio di non avere la possibilità di smaltire sul mercato tutti i quantitativi in stock (a meno di sven-dite svantaggiose);

d'altro lato scorte troppo basse danno luogo a:

- frequenti situazioni di rottura di stock (mancanza degli articoli richiesti) con conseguente perdita di guadagno
- numerosi solleciti ed intralci al normale flusso lavorativo
- perdita dell'immagine commerciale dell'azienda presso la clientela, che può portare anche alla perdita di significative percentuali di mercato.

Normalmente nella vita aziendale la ricerca di un equilibrio fra queste opposte esigenze conduce a vivaci scontri fra i « finanziari » preoccupati per gli immobilizzi ed i « commerciali » che ovviamente vorrebbero una pronta evadibilità per ogni ordine.

Inoltre tutto è reso più complesso in periodo di congiuntura caratterizzata da forti variazioni di prezzo quando considerazioni speculative prevalgono distorcendo la normale politica di gestione.

2. Problemi della gestione scorte

Molto efficacemente si usa spesso riassumere il problema fondamentale della gestione delle scorte nei due semplici interrogativi:

- quando ordinare?
- quanto ordinare?

Per dare risposte valide, a livello dei singoli articoli del magazzino che si intende gestire, è fondamentale impostare un sistema efficiente di previsione delle vendite.

Esso infatti costituisce il punto focale di tutto il problema perché ottenute previsioni precise si hanno i dati essenziali per una buona gestione degli stock.

2.1. - Previsioni di Vendita

I metodi automatici di previsione si possono raggruppare sostanzialmente in due grandi categorie:

— *Metodi regressivi*: essi cercano correlazioni fra le

vendite del prodotto in esame ed altre variabili che influenzano le vendite stesse.

Per costruire un buon modello di questo tipo è necessario uno studio dettagliato dei fenomeni che possono influire sulle vendite. Attraverso programmi di regressione si può analizzare il grado di correlazione esistente tra le varie « cause » che si ipotizza influenzino le vendite e le vendite stesse. In questo modo è possibile scartare dal modello quelle variabili (« cause ») che risultano poco e per nulla correlate con le vendite.

Anche se dal punto di vista « software » ci si può servire di programmi in genere forniti dai costruttori, la messa a punto di modelli di questo tipo richiede un notevole lavoro di analisi giustificabile solo per pochi articoli molto importanti o per linee omogenee di prodotti.

D'altra parte è opportuno sottolineare che questo tipo di approccio è l'unico con il quale si può sperare di avere risultati soddisfacenti per previsioni a medio-lungo termine.

— *Metodi autoregressivi*: una drastica semplificazione nella costruzione del modello previsionale interviene quando si rinunci alla ricerca delle variabili che influenzano le vendite del prodotto considerato, e ci si limiti a previsioni ottenute estrapolando i dati storici delle vendite passate.

Nella ipotesi che le cause influenzanti il fenomeno vendita rimangano dello stesso tipo e non subiscano grosse modifiche, utili risultati possono venire da uno studio dell'andamento delle vendite nel tempo, che si traduce in una analisi delle componenti cicliche, delle tendenze di fondo (trend) e della variabilità della serie di dati storici. La curva delle vendite così analizzata viene poi proiettata nel futuro.

Questi criteri, per la loro stessa natura, sono utilizzabili solo per previsioni a breve termine in quanto le fluttuazioni del mercato (dovute a lanci di nuovi prodotti, alla concorrenza, a campagne pubblicitarie, ecc...) rendono troppo aleatoria la estrapolazione dei dati passati di vendita su archi di medio o lungo termine.

Il mancato rispetto di questi limiti di validità è stata una delle principali cause dei risultati poco validi lamentati da alcuni utilizzatori di queste tecniche.

Un secondo errore da evitare è l'applicazione di questi metodi ad articoli soggetti a fenomeni di « moda » che possono dare alle curve di vendita impennate o cadute del tutto scorrelate con l'andamento precedente.

A volte andamenti strani od « erratici » sono caratteristici anche di prodotti che a priori sembrerebbe dover avere un comportamento regolare.

Questo è sempre più vero per articoli soggetti a fenomeni commerciali molto variabili e per i quali spesso risulta impossibile qualsiasi procedimento automatico. In questi casi la stima delle previsioni è unicamente demandata al marketing.

Negli altri casi, se usati con le dovute cautele, i

metodi autoregressivi possono dare risultati molto soddisfacenti, soprattutto in rapporto alla semplicità delle tecniche usate.

D'altra parte quando si tratta un gran numero di articoli è impensabile spendere per ognuno il tempo di analisi richiesto dai modelli di regressione, per cui in pratica gli unici metodi automatici proponibili sono quelli autoregressivi, che ora esamineremo con maggiori dettagli.

3. Caratteristiche e tipi di modelli autoregressivi

3.1. Medie semplici

Fra questi metodi si possono includere alcune delle tecniche più elementari:

- *media aritmetica dei dati storici*

Si valuta un consumo medio di ogni articolo semplicemente come media aritmetica di tutti i dati storici. Il valore ottenuto viene estrapolato sui mesi futuri per ottenere una previsione.

- *media aritmetica mobile dei dati storici*

Simile al precedente, ma con media limitata solo agli ultimi n valori storici (es. $n = 6$). Questo permette di limitare il numero dei dati storici necessari con vantaggio per le dimensioni degli archivi.

- *media mobile pesata dei dati storici*

In questo caso anziché eseguire una semplice media aritmetica si attribuiscono *pesi* diversi ai dati storici, normalmente maggiori ai dati più recenti e via via decrescenti per i dati più vecchi.

E' infatti ragionevole ritenere che le situazioni del mercato evolvano in modo tale da rendere meno significativi i dati più lontani rispetto ai più recenti. Fra i vari tipi di medie pesate, ha avuto molto successo il metodo di smorzamento esponenziale (exponential smoothing nella terminologia anglosassone) che illustreremo.

3.2. Smorzamento esponenziale

Se indichiamo con v_t il venduto nel periodo t (il periodo può essere il mese, la settimana od altra base di tempo), si calcola periodicamente una media (detta stima) $\bar{v}_{t+1}$ pari al valore del venduto previsto per il periodo $t+1$ usando l'espressione:

$$\bar{v}_{t+1} = \bar{v}_t + \alpha (v_t - \bar{v}_t) \quad (1)$$

dove:

$\bar{v}_t$ = stima calcolata al mese precedente

α = parametro di valore compreso fra 0 ed 1 (sul suo significato si discuterà in seguito).

E' facile riconoscere che la stima così calcolata risulta essere una forma particolare di media pesata delle vendite. Infatti per la natura iterativa della (1) il valore della stima $\bar{v}_t$ può a sua volta essere espresso in funzione di $\bar{v}_{t-1}$ e di v_{t-1} .

A sua volta $\bar{v}_{t-1}$ è esprimibile in funzione dei dati al periodo $t-2$ e così via.

Iterando il procedimento si riconosce che:

$$\bar{v}_{t+1} = \alpha \cdot v_t + \alpha (1 - \alpha) v_{t-1} + \alpha (1 - \alpha)^2 v_{t-2} + \dots + \alpha (1 - \alpha)^i v_{t-i} + \dots \quad (2)$$

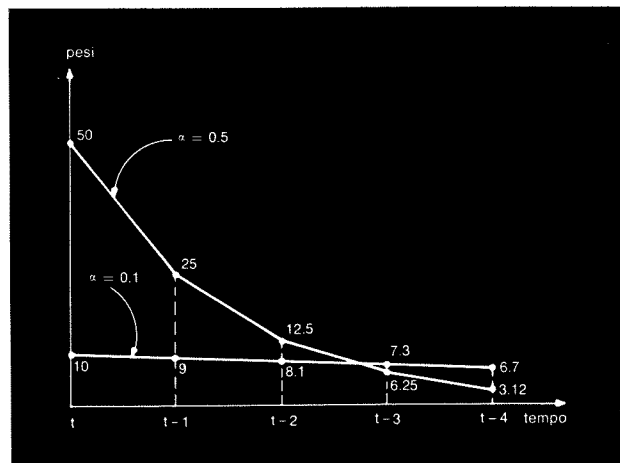


Fig. 1 - Previsioni di vendita secondo il modello di smorzamento esponenziale.

Dalla (2) risulta chiaro che la stima $\bar{v}_{t+1}$ è espressa da una media dei dati di vendita $v_t, v_{t-1}, \dots, v_{t-i}, \dots$. I pesi hanno valore:

$$\alpha, \alpha(1-\alpha), \alpha(1-\alpha)^2, \dots, \alpha(1-\alpha)^i, \dots \quad (3)$$

Si riconosce facilmente che i pesi costituiscono una serie geometrica di ragione $(1-\alpha)$ e per $0 < \alpha < 1$ la loro somma converge ad 1 come richiesto per la normalizzazione della media.

Spesso questo modello è usato per calcoli di previsioni. La sua praticità risiede nel poter ricalcolare la nuova stima (cioè la nuova previsione di consumo) tramite la semplice espressione (1) che richiede solo due dati:

- la vecchia stima $\bar{v}_t$
- l'ultimo dato del venduto v_t

al contrario delle medie illustrate precedentemente che richiedevano sempre n dati storici (questo significa risparmio nella dimensione degli archivi).

3.3. La scelta del parametro α

Nell'utilizzo di questo modello il problema chiave è la scelta del valore di α . Per vederne la ragione, esaminiamo la (3) che esprime la successione dei pesi attribuiti ai dati storici.

Consideriamo due casi: $\alpha = 0,5$ ed $\alpha = 0,1$.

Esprimendo i pesi in valore percentuale, le due successioni che risultano sono:

$$\begin{array}{cccccc} 50\% & 25\% & 12,5\% & 6,25\% & 3,12\% & \dots \\ 10\% & 9\% & 8,1\% & 7,30\% & 6,70\% & \dots \end{array} \quad \begin{array}{l} (\alpha = 0,5) \\ (\alpha = 0,1) \end{array}$$

dove il primo dato si riferisce al valore più recente conosciuto, ed i successivi ai valori via via più vecchi. Gli andamenti sono riportati in fig. 1.

Si nota che nel primo caso la stima dei dati di vendita è in pratica legata solo ai valori più recenti (da soli i primi tre dati concorrono per l'87,5% a formare la stima).

Dunque appare chiaro che un modello con α elevato tende a « dimenticare » in fretta i dati più vecchi ed a seguire con prontezza gli andamenti delle vendite. Al contrario nel caso di $\alpha = 0,1$ i primi tre valori concorrono solo per il 27,1% e quindi la storia più

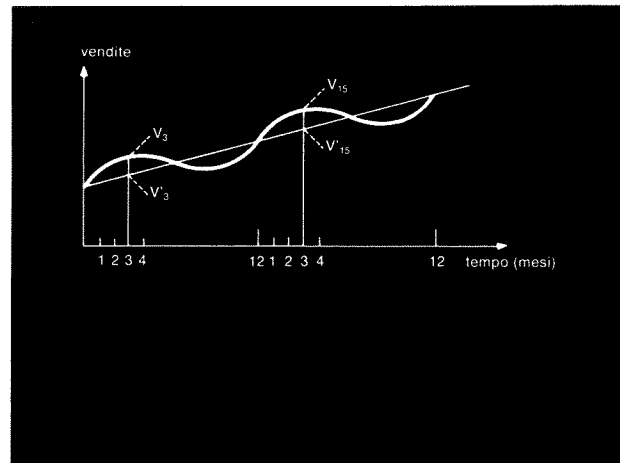


Fig. 2 - Valutazione dei coefficienti stagionali di vendita.

lontana conserva una influenza notevole. Ne risulta che per valori di α piccoli il modello reagisce con lentezza a variazioni nella serie storica « filtrandoli » notevolmente.

Dunque α piccoli sono validi quando l'andamento dei dati è turbato da brusche anomalie saltuarie (picchi isolati dovuti a vendite casualmente molto alte o molto basse) che è bene filtrare per non perturbare le previsioni.

Riassumendo non esiste un valore di α ottimo in generale. La scelta viene fatta, con metodi automatici, articolo per articolo e deve essere rivista costantemente nel tempo.

Nei sistemi più moderni il parametro è autoregolato in funzione degli scarti osservati fra previsioni e dati effettivi. In questo modo il modello si autocorregge automaticamente (1).

3.4. Modelli autoregressivi complessi

I modelli di tipo exponential smoothing corentemente usati sono più complessi di quello illustrato, in quanto la estrapolazione dei dati storici è fatta tenendo conto di altre due componenti fondamentali: — la presenza di fluttuazioni cicliche ripetitive (fenomeni di stagionalità)

— una tendenza di fondo (trend) a crescere od a decrescere della serie storica (caso di articoli che stanno affermandosi sul mercato o che sono in via di estinzione).

E' ovvio che un modello che ignorasse queste componenti basandosi su una semplice media pesata dei dati passati sarebbe soggetto al rischio di errori sistematici.

3.5. Valutazione dei coefficienti stagionali

Il calcolo dei coefficienti stagionali presuppone la conoscenza di almeno due anni di storia delle vendite. L'analisi potrebbe essere fatta per ogni singolo articolo per ricavare i suoi coefficienti stagionali, ma in

(1) Questa tecnica è utilizzata nel software applicativo AIMS (Autoadaptive Inventory Management System) e DIMS (Distribution Inventory Management System) su vari modelli di elaboratori Honeywell.

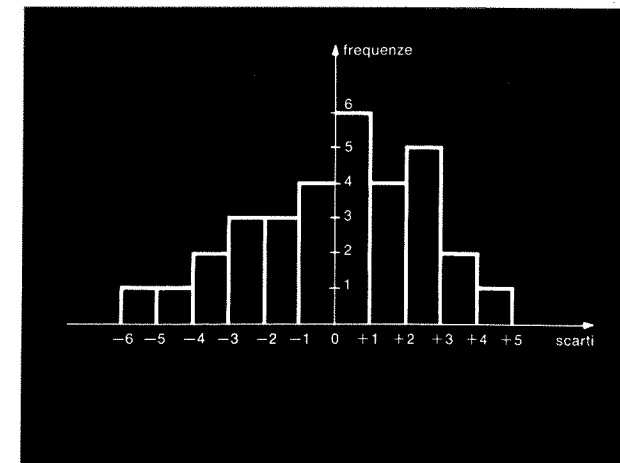


Fig. 3 - Distribuzione degli scarti riscontrati.

genere ci si limita a ricercare coefficienti per classi di prodotti ad andamento simile.

I dati di vendita sono solitamente su base mensile (o per gruppi di 4 settimane) e devono comprendere le sole vendite ordinarie, cioè essere depurati di eventuali vendite straordinarie.

Per chiarire quest'ultimo punto pensiamo di considerare la storia delle vendite di un pezzo di ricambio che abbia la serie:

30, 25, 12, 328, 22, 33, 18, ...

Colpisce immediatamente la sproporzione del 4° dato rispetto alla media degli altri. Risalendo alle cause può risultare che il 328 in questione sia la somma di 28 pezzi usciti al « minuto » più una partita di 300 pezzi ordinata « una tantum » da un grossissimo cliente.

Ora è necessario tralasciare i 300 pezzi ed occuparsi solo delle vendite ordinarie per non falsare il calcolo della stagionalità.

Da questo può derivare un primo suggerimento per utilizzatori che vogliano prepararsi ai metodi previsionali automatici: raccogliere i dati storici delle vendite, tenendo separate eventuali vendite eccezionali, come nell'esempio.

Quasi tutti i metodi usati in pratica procedono interpolando con una retta la serie di dati del venduto, e quindi ricavando i coefficienti stagionali come rapporto fra dati effettivi ed i valori che si hanno in corrispondenza sulla retta.

In fig. 2 il coefficiente stagionale di Marzo viene calcolato come:

$$CS_3 = \frac{v_3 + v_{15}}{v'_3 + v'_{15}}$$

E' possibile il calcolo dei coefficienti relativi ad una famiglia di articoli, con controllo della ipotesi che tutti gli articoli abbiano stagionalità simili.

Tecniche avanzate del tipo « cluster analysis », le cui applicazioni a questo settore sono agli inizi, consentono di individuare automaticamente gruppi di articoli con stagionalità simile, proponendo per essi una sola serie di coefficienti stagionali.

In tutti i casi è bene sottolineare che la esperienza acquisita in questo tipo di applicazioni consiglia sempre un esame critico dei risultati da parte dei responsabili del marketing per evitare abbagli grossolani.

3.6. Il modello completo

Definiti i coefficienti stagionali, il modello ricava una serie di dati storici « destagionalizzati », cioè depurati delle oscillazioni ripetitive.

Su questi dati viene calcolata la pendenza media di una retta interpolante (2); la previsione futura è quindi costruita a partire dalla media mobile dei dati passati corretta del termine di trend così calcolato.

Ottenuta la previsione sui dati destagionalizzati è quindi sufficiente reintrodurre i coefficienti stagionali per ottenere la previsione completa.

In formule, la previsione di vendita $\bar{v}_{t+1}$ relativa al periodo $t+1$ viene espressa come:

$$\bar{v}_{t+1} = (\bar{v}_{t+1}^* + T_{t+1}) \cdot C_{t+1} \quad (4)$$

dove:

C_{t+1} è il coefficiente stagionale relativo al periodo $t+1$
 T_{t+1} è il termine di trend calcolato all'inizio del periodo $t+1$

$\bar{v}_{t+1}^*$ è la « stima » dei dati di vendita ottenuta dal dato del venduto dell'ultimo periodo v_t (destagionalizzato dividendolo per C_t) secondo la classica espressione dell'exponential smoothing:

$$\bar{v}_{t+1}^* = \bar{v}_t^* + \alpha \left(\frac{v_t}{C_t} - \bar{v}_t^* \right) \quad (5)$$

dove α è ancora il fondamentale parametro del cui significato si è discusso.

Il trend T è spesso calcolato secondo formule ricorrenti di tipo « exponential smoothing », od attraverso interpolazioni dei dati storici.

Questo tipo di modello (che in pratica è il più diffuso) corrisponde al metodo noto in letteratura come « exponential smoothing del II ordine » (trend di tipo lineare) con coefficienti di stagionalità moltiplicativi (cioè oscillazioni stagionali di ampiezza proporzionale al valore medio della serie).

E' possibile rinunciare al termine di trend, utilizzando per le previsioni la (4) semplificata come:

$$\bar{v}_{t+1} = \bar{v}_{t+1}^* \cdot C_{t+1}$$

Questo modello, noto come « exponential smoothing del I ordine », fornisce una estrapolazione orizzontale della stima dei dati di vendita e può essere prudente utilizzarlo nei casi in cui il calcolo del trend è poco affidabile (serie instabile).

4. Distribuzione statistica delle previsioni

E' chiaro che il fenomeno degli ordini futuri che stiamo cercando di prevedere è per sua natura una variabile aleatoria, descrivibile solo in termini statistici.

(2) Talvolta si considerano in letteratura interpolazioni paraboliche o polinomiali di grado anche più alto. Nella pratica espressioni del genere non sono quasi mai usate.

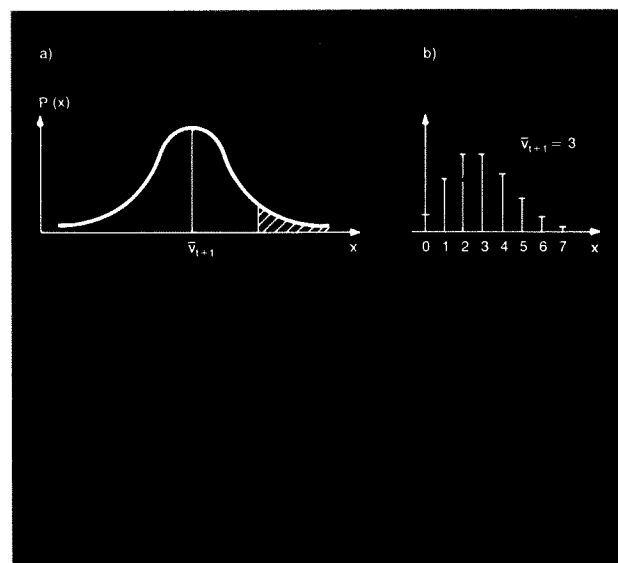


Fig. 4 - Distribuzione statistica degli scarti: a) distribuzione di Gauss, b) distribuzione di Poisson.

Supponiamo che dalla interpolazione dei dati storici secondo le tecniche descritte si sia ottenuto un valore che chiameremo $\bar{v}_{t+1}$ che rappresenta la previsione su di un periodo futuro.

La prima domanda che si pone è: quali scarti è ragionevole attendersi rispetto a $\bar{v}_{t+1}$?

Infatti nella ipotesi di comportamento stazionario della serie storica studiata potremo supporre che rispetto a $\bar{v}_{t+1}$ scarti positivi e scarti negativi siano egualmente probabili, ma è importante avere indicazioni sulla loro entità.

In termini più esatti si può dire che per definire il fenomeno, va precisata la distribuzione statistica di probabilità che lo descrive.

Considerando i prodotti uno ad uno si può pensare di costruire in modo empirico le curve che riportano le frequenze degli scarti rispetto ai valori attesi.

Le spezzate così ottenute (fig. 3) approssimano, al crescere del numero delle osservazioni, la distribuzione di probabilità attorno al valore previsto. Osserviamo però che costruire curve singole di questo tipo porterebbe in generale a distribuzioni diverse e variabili nel tempo, con enorme complessità di trattazione.

Per trattare automaticamente migliaia di articoli è opportuno ricondursi a pochi tipi di distribuzione che possano mediamente approssimare i diversi tipi di comportamento.

Dalla statistica si ha che la curva di Gauss può essere scelta come la approssimazione più ragionevole quando la movimentazione dei prodotti è espressa da numeri abbastanza alti (fast moving items).

In molti modelli matematici di gestione scorte questo tipo di distribuzione è assunto come valido indiscriminatamente per tutti i prodotti.

In realtà per gli articoli di bassa movimentazione la gaussiana risulta assolutamente inadeguata, mentre la distribuzione di Poisson (fig. 4) meglio approssima la distribuzione effettiva.

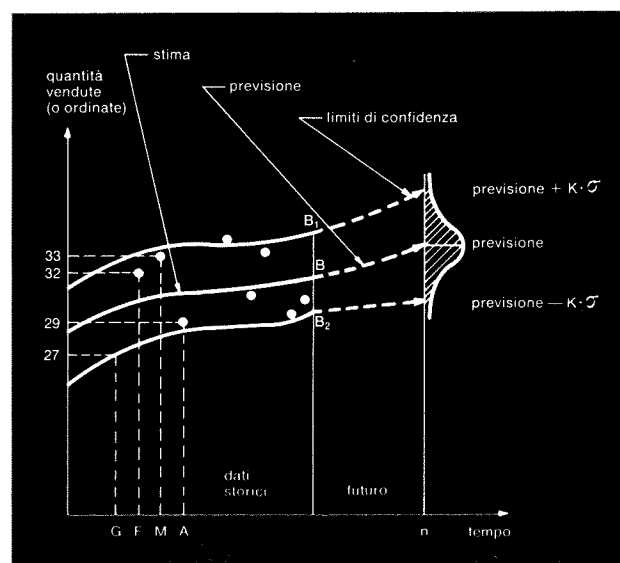


Fig. 5 - Filtraggio dei dati di vendita.

I più evoluti modelli di gestione (tipo l'AIMS) provvedono a classificare automaticamente i prodotti in famiglie di alta o bassa movimentazione così da utilizzare per essi uno o l'altro tipo di distribuzione per il calcolo della dispersione attorno al valore medio. La classificazione è rivista dinamicamente in modo da passare i prodotti da una classe all'altra, quando la loro movimentazione lo giustifichi.

Definita la natura della curva, un solo parametro, la varianza, è sufficiente a descrivere univocamente il grado di dispersione attorno al valore medio.

La varianza è calcolata per gli articoli ad alta movimentazione a partire dal valore medio degli scarti fra previsioni e valori effettivi moltiplicato per un coefficiente fisso.

Nella distribuzione di Poisson la varianza risulta pari alla radice quadrata del valore medio $\bar{v}_{t+1}$.

E' interessante notare che per il calcolo delle varianze non si richiedano nuovi dati di input essendo sufficienti i valori della serie storica necessari per il calcolo dei valori medi pesati.

5. Filtro dei dati di vendita

Calcolata la previsione e la relativa varianza, è interessante evidenziare i prodotti per i quali i dati effettivi escono da un « intervallo di previsione » costruito attorno a $\bar{v}_{t+1}$.

Questo permette di fornire un tabulato per eccezioni contenente i soli articoli che hanno avuto come ultimo dato di vendita un valore molto lontano da quello atteso.

La segnalazione è importante soprattutto per valori eccessivamente alti in quanto gli articoli corrispondenti sono sicuri candidati alla rottura di stock.

In questi casi, oltre ad un controllo sulla ragione della vendita anomala, la segnalazione può permettere azioni tempestive (ad es. sollecito telefonico di ordini aperti).

L'ampiezza del « filtro » con il quale si vagliano

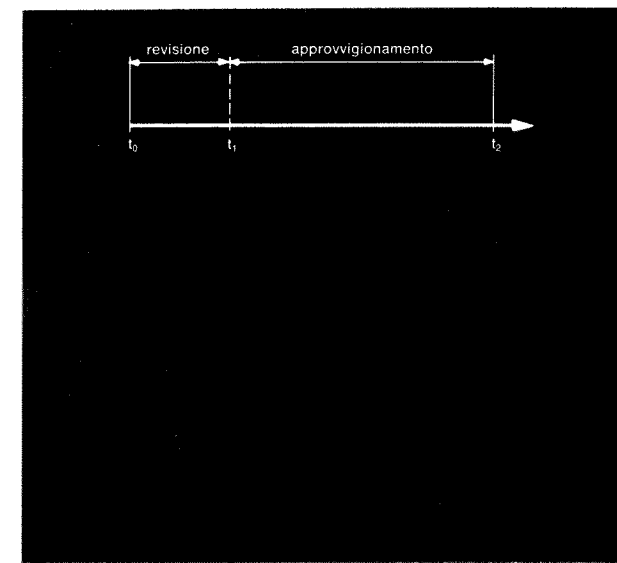


Fig. 6 - Tempi di revisione e di approvvigionamento.

i dati viene di solito fissata empiricamente così da intercettare solo gli scarti ritenuti significativi come ordini di grandezza (vedi fig. 5).

In teoria un filtro abbastanza ampio sarebbe ottenibile sommando e sottraendo attorno a $\bar{v}_{t+1}$ un intervallo pari a due o tre volte la varianza. Ad esempio nel caso teorico di distribuzione gaussiana, un intervallo pari a tre volte la varianza dovrebbe contenere il dato reale nel 99,7% dei casi. In pratica per l'influenza di fenomeni particolari si osserva che questa percentuale è di solito più bassa e, per classi di articoli particolarmente anomale, a volte sono usati anche filtri di ampiezza maggiore.

6. Tempi di approvvigionamento

Le previsioni di cui si è trattato erano relative ad un arco di tempo pari ad un periodo di base (settimana, mese, ecc.). Ai fini di una applicazione di gestione delle scorte interessa spingere la previsione su un periodo di tempo (detto « lead time » nella terminologia anglosassone) pari alla somma di:

— un tempo di consegna, inteso come intervallo complessivo di tempo dalla decisione del riordino al momento della accettazione della merce a magazzino;

— un tempo di revisione, pari all'intervallo di tempo intercorrente fra due successive elaborazioni che rivedono la situazione dello stock.

Infatti supponiamo che una elaborazione sia fatta al tempo t_0 (fig. 6); dopo di essa si dovrà attendere almeno fino a t_1 (dopo un tempo di revisione, alla prossima elaborazione) per emettere una proposta di riordino con consegna non prima di t_2 . Per questa ragione già al tempo t_0 è necessario calcolare il fabbisogno previsto fino a t_2 .

Naturalmente i tempi di approvvigionamento variano in generale da articolo ad articolo ed, inoltre, cambiano in funzione di tante variabili (situazione del mercato, scioperi presso il produttore o nei trasporti,

ecc.). Per questo spesso risulta difficile disporre di loro valori aggiornati.

Quello che un sistema automatizzato può fare è di confrontare costantemente i tempi effettivi (valutabili all'accettazione delle merci in arrivo per confronto con le date degli ordini) aggiornando i dati di archivio.

E' ovvio inoltre che variazioni conosciute nei termini di consegne vanno comunicate al sistema al più presto, senza attendere il loro effettivo riscontro.

7. Stock di sicurezza e punto di riordino

Calcolata la previsione di consumo su di un arco di tempo pari ad un lead time, la gestione dello stock si riconduce ad un semplice confronto fra la disponibilità del prodotto ed il fabbisogno preventivato.

Per disponibilità intendiamo: esistenza fisica + quantità ordinata — quantità impegnata.

Dal confronto scaturisce la necessità di riordinare o meno il prodotto.

In pratica per cautelarsi adeguatamente verso probabili rotture di stock si somma alla previsione $\bar{v}_{t+1}$ una « scorta di sicurezza » v_s che permetta di far fronte, entro certi limiti, ad eventuali picchi di consumo.

La determinazione di v_s è molto delicata perché un suo errato dimensionamento può portare da un lato ad immobilizzi eccessivi e dall'altro a frequenti rotture di stock.

In una gestione automatizzata la dimensione dello stock di sicurezza è ricalcolata dinamicamente per ogni prodotto in funzione di una serie di parametri che sono:

— la varianza della distribuzione di probabilità relativa al prodotto; è intuitivo che lo stock di sicurezza crescerà al crescere di questo valore;

— il livello di servizio che si desidera avere per il prodotto. Ricordiamo che per un magazzino con consegna a pronto la definizione di livello di servizio è data dal rapporto:

$$\frac{n. \text{ articoli consegnati}}{n. \text{ articoli richiesti}}$$

— la dimensione dei lotti di acquisto. Di questo parametro si parlerà nel par. 8.

Per intuire la sua relazione con lo stock di sicurezza, basta osservare che, a parità di livello di servizio atteso, tanto più alti sono i lotti d'acquisto, tanto minore può essere v_s .

Riassumendo, possiamo osservare che la somma della previsione di vendita e della scorta di sicurezza fornisce un parametro che, per rifarci alla terminologia corrente, può essere definito il « punto di riordino ».

Il vantaggio della tecnica descritta rispetto al metodo di lavoro correntemente in uso sta nel continuo aggiornamento del punto di riordino così calcolato che risulta sempre calibrato all'effettivo comportamento del prodotto.

8. Lotti d'acquisto

Dal confronto descritto fra punto di riordino e disponibilità il sistema verifica se è il momento di riordinare: questo risponde alla domanda sul *quando* ordinare.

Resta da specificare *quanto* conviene richiedere al fornitore.

Sulla determinazione dei lotti di acquisto giocano molti fattori economici, tecnici, logistici.

Una trattazione ormai classica del problema ricerca il lotto come compromesso ottimale fra varie componenti di costo che possiamo elencare:

- costi di emissione e gestione ordini, tanto più alti quanto più numerosi sono gli ordini stessi. E' ovvio che per diminuire l'influenza di questo termine converrebbe emettere pochi ordini per grossi quantitativi;
- costi di immagazzinaggio ed immobilizzo capitali proporzionali alle giacenze medie, per ridurre i quali sono invece opportuni piccoli lotti;
- risparmi conseguenti a sconti per quantitativi elevati.

Se altre considerazioni non sono prevalenti è possibile calcolare il lotto ricercando un compromesso ottimale fra le voci elencate, ma spesso problemi commerciali o finanziari pongono limiti pratici che superano ogni calcolo teorico. Ad esempio del tutto particolare è in ogni caso la trattazione di articoli che abbiano una forte deperibilità (shelf life items).

Altre volte i fornitori concedono interessanti condizioni legate non a singoli prodotti, ma al totale dell'ordine.

In questo caso l'acquisto più vantaggioso può risultare quello in cui l'ordine ad un fornitore contiene un valore complessivo pari almeno ad una certa soglia (espressa in una certa unità di misura: Lire, Kg., litri, ecc.).

In pratica per questi casi sono usate delle procedure automatiche che provvedono a formulare una proposta d'ordine che complessivamente raggiunge il totale richiesto, distribuendo in modo omogeneo sulle varie righe d'ordine le eventuali eccedenze necessarie.

Questo è calcolato tenendo presente sia le previsioni di consumo che le disponibilità dei vari prodotti, così da non squilibrare la gestione creando stock eccessivi solo per alcuni di essi.

9. Le procedure EDP

9.1. Le fasi di lavoro

Per tradurre i concetti esposti in procedure funzionanti sull'elaboratore occorre creare una serie di programmi che possiamo suddividere in quattro grandi categorie in funzione della loro frequenza di utilizzazione.

1) *Procedure di inizializzazione* - Queste fasi analizzano la storia delle vendite ricavando i coefficienti stagionali ed inizializzando tutte le formule di tipo ricorrente. Di solito la frequenza è annuale.

2) *Procedure di ottimizzazione* - Fra queste comprendiamo una serie di fasi aperiodiche realizzate per la

ricerca delle condizioni ottimali di funzionamento della gestione sia per quanto riguarda il livello di servizio sia per i lotti.

Nei vari modelli utilizzati, il lotto di acquisto è funzione di diversi parametri fra i quali quasi sempre il costo del prodotto stesso, la sua previsione di vendita annuale ed altre variabili legate ai costi di gestione.

Non c'è una regola fissa per aggiornare i lotti, solitamente lo si fa su base semestrale od annuale salvo revisioni straordinarie in certe occasioni precise (ad es. cambio del listino).

3) *Aggiornamento del Punto di Riordino* - Queste fasi ricalcolano le previsioni di vendita e gli stock di sicurezza elaborando le formule ricorrenti descritte nel par. 3. La periodicità è quella della rilevazione dei dati di vendita, ma quasi sempre mensile. In questa fase si emette il tabulato originato dal filtro dei dati di vendita.

4) *Emissione proposte di riordino* - Con frequenza almeno pari alle fasi di tipo 3, ma generalmente maggiore, viene elaborato un programma che confronta il disponibile con il punto di riordino ed emette le proposte di acquisto. In queste fasi sono spesso emesse anche segnalazioni di allarmi per articoli la cui esistenza fisica è scesa sotto certi valori critici.

9.2. Caratteristiche tecniche

Solitamente tutte queste fasi non necessitano di dimensioni di memoria particolarmente elevate, in quanto le maggiori difficoltà sono di tipo concettuale per la complessità degli aspetti matematici e statistici piuttosto che per la mole dei calcoli.

Per alcuni modelli matematici di previsione questo è vero a meno delle fasi di inizializzazione, che possono richiedere l'utilizzo di sistemi medio-grandi. Per esempio, una ottimizzazione molto spinta si ottiene usando tutte le routines AIMS su sistemi della serie G-100 a partire da soli 12 K di memoria.

Per quanto riguarda la struttura delle periferiche la sequenza illustrata non ha particolari requisiti perché ogni articolo è gestito singolarmente (si tratta cioè di elaborazioni sequenziali).

In termini di impegno di memoria di massa per automatizzare la gestione di un magazzino, del quale già sia meccanizzata la contabilità delle entrate e delle uscite, è sufficiente aggiungere al record prodotto sull'archivio articoli un certo numero di campi necessari per i calcoli matematico-statistici descritti.

L'impegno che ne risulta non supera i 100 caratteri per articolo gestito.

Un salto di complessità nella struttura degli archivi si ha invece nei casi di riordini multipli (joint ordering) per i quali è necessario gestire un archivio prodotti su disco con opportuni legami fra i prodotti relativi ad uno stesso fornitore.

Circa i tempi di esecuzione le fasi più pesanti sono certamente quelle di inizializzazione, che però sono poco frequenti.

L'utilizzo di routine AIMS permette di contenere

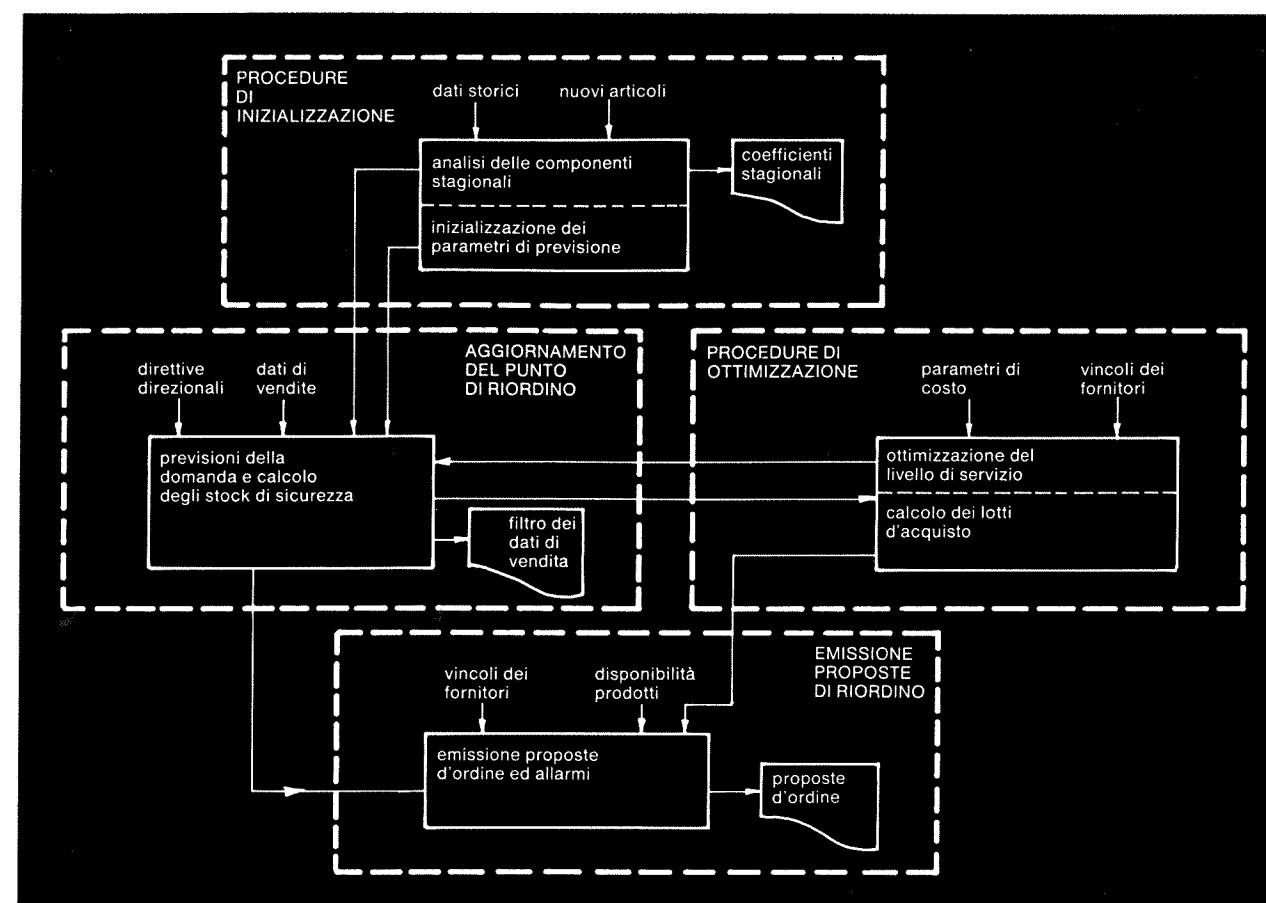


Fig. 7 - Le procedure EDP.

questo tempo in modo notevole, tanto che la velocità media di inizializzazione per un elaboratore della classe G-100 è risultata in applicazioni effettive di 1.000 articoli/ora circa per storie di 24 mesi. Presso le stesse installazioni le fasi mensili di aggiornamento del punto di riordino elaborano circa 8-10.000 articoli/ora, mentre la fase settimanale di emissione proposte d'acquisto è ancora più veloce in quanto si limita a calcoli banali.

In questo caso non è il tempo di elaborazione nella unità centrale, ma quello delle periferiche a determinare il tempo di esecuzione.

In generale si può osservare che (in assenza di joint ordering) tutte le fasi hanno durata direttamente proporzionale al numero degli articoli gestiti.

In fig. 7 sono riportate, a grandi linee, le funzioni e le interconnessioni fra le diverse procedure.

10. Considerazioni conclusive

Il tipo di gestione che abbiamo descritto si complica quando i magazzini sono più di uno e, come spesso accade in pratica, costituiscono una rete complessa. Questo è il caso, ad esempio, di aziende articolate su filiali dotate di depositi con importanza solo regionale

e di magazzini principali (eventualmente presso le fabbriche).

In questo caso il problema che si pone è la gestione dell'intero sistema complesso mirante a minimizzare da un lato le scorte totali, dall'altro i trasferimenti interni soprattutto fra magazzini secondari.

Si devono così studiare soluzioni che legano strettamente il problema della gestione stock (ancora affrontabile con tecniche del tipo descritto) a quello del trattamento e della gestione ordini spingendo verso soluzioni avanzate che, dal punto di vista EDP, prevedono sempre più spesso un largo uso di terminali per un rapido scambio di informazioni e decisioni fra periferia e centro.

11. Bibliografia

- [1] R.G. BROWN: *Statistical forecasting for inventory control*; Mc-Graw Hill, New York.
- [2] J.F. MAGEE: *Production planning and inventory control*; Mc-Graw Hill, New York.
- [3] A.F. VEINOTT: *The status of mathematical inventory theory*; Management science, vol. 12, 745-777 (1966).
- [4] *Il sistema AIMS per la gestione delle scorte* - Rapporto Honeywell.
- [5] *Descrizione funzionale DIMS - Serie 60 Livello 62* - Manuale Honeywell.

Che cosa possono fare e che cosa non possono fare le macchine

J. NIEVERGELT, J.C. FARRAR

University of Illinois (USA)
Department of Computer Science

1. Introduzione

Quando ci si chiese, per la prima volta, che cosa possano fare e che cosa non possano fare le macchine, ciò non avvenne certamente in relazione ai calcolatori. Sembra che l'umanità sia stata sempre curiosa di trovare i limiti a cui può giungere la tecnologia di un'epoca. In particolare, una grande attrazione è sempre stata suscitata dall'idea di costruire macchine che possano, sotto qualche aspetto, simulare il comportamento dell'uomo. Vi è sempre stato qualche inventore che ha sognato, ha progettato, e qualche volta ha costruito macchine, quali un guerriero a vapore o uno scrivano automatico ad orologeria.

Tuttavia, per effetto del grande aumento di velocità di calcolo e di complessità delle macchine, raggiunto negli ultimi decenni, e del radicale miglioramento nella nostra comprensione del modo di progettare algoritmi tali da fare sì che la macchina si comporti come vogliamo, tale domanda ha assunto importanza maggiore come non mai prima d'ora. Ed inoltre, data la sempre maggiore diffusione dei calcolatori e delle macchine comandate da essi nella nostra società è importante che il pubblico, in generale, comprenda almeno in parte le possibilità e le limitazioni di tali macchine.

La domanda: « Che cosa può fare, e che cosa non può fare una macchina? » non ammette una risposta netta, se non altro perché le parole « può fare » sono così vaghe da permettere diverse interpretazioni. Negli ultimi quarant'anni, però, alcune domande precise riferentisi ad essa sono state formulate, e alcune risposte parziali sono state ottenute. La conoscenza e la comprensione di esse può servire ad evitare confusione circa le capacità delle macchine, e a dissipare certi miti largamente diffusi.

2. Una macchina può pensare?

Questa è forse una delle domande più recenti e stimolanti che si possano fare riguardo a una macchina.

Fino all'avvento dei calcolatori digitali, le macchine erano progettate quasi esclusivamente per effettuare operazioni puramente meccaniche, e davano perciò scarso motivo di riflessione sulle loro capacità intellettuali.

Vi sono stati tuttavia esempi antichi che poterono far sorgere, nelle menti di molte persone, il problema della intelligenza delle macchine. Vi fu, al riguardo, il caso del falso giocatore automatico di scacchi costruito dal Barone von Kempelen intorno al 1800, che fu esposto al pubblico per molti anni, e di cui si dice che abbia vinto tutte le partite giocate. Entro la macchina era così abilmente celata una persona che, benché fosse permesso agli spettatori di ispezionare, successivamente, le diverse parti interne della macchina, egli, spostandosi dall'una all'altra parte, riuscì sempre ad evitare di essere scoperto.

Oggi esistono diversi programmi di gioco degli scacchi, che giocano al livello dei giocatori di torneo delle qualifiche inferiori, e che batterebbero la maggioranza dei giocatori occasionali. V'è anche un programma che ha raggiunto il livello di « maestro ».

I più saranno disposti ad ammettere che questa prestazione è un risultato di natura intellettuale, non banale. Tuttavia, la discussione sulla capacità delle macchine di adempiere compiti di natura intellettuale è stata spesso oscurata da argomenti emotivi ed irrazionali.

3. La prova di Turing

La questione delle macchine pensanti fu discussa nel 1950 dal famoso logico inglese A.M. Turing, e, ancor oggi, la sua analisi è una delle migliori che si possano leggere da chi voglia chiarire le proprie idee al riguardo. Turing vede chiaramente che la domanda è di natura emotiva, tale da prestarsi a interpretazioni soggettive su cosa significhi « pensare », e ben presto la respinge come troppo vaga per avere una risposta significativa. Infatti, egli indica come, per

certe interpretazioni della domanda: « la macchina può pensare? » la risposta non possa ottenersi con ragionamenti logici, ma solo come materia di fede: mentre, per altre, la risposta è una conclusione anticipata. (Certi argomenti circa le macchine pensanti sono basate sulla assunzione inespressa che il pensare sia, per definizione, ciò che le persone fanno: da cui si deduce che le macchine non possono pensare, a meno che si voglia prendere in considerazione la possibilità che le macchine siano persone, o che le persone siano macchine: il che non intendiamo fare.)

Turing ha posto una nuova questione, al posto della domanda originale: un modello secondo cui il « pensiero » delle macchine può essere discusso in termini puramente tecnici. Esso è noto da allora come la « prova (test) di Turing » ed è descritta come segue. Si immagini un calcolatore C, e una persona M rinchiusi in una stanza, e collegati ognuno ad una telescrivente in un'altra stanza, dove c'è un interrogatore umano, I. Il compito di I è di determinare quale telescrivente è collegata a C, e quale ad M, « battendo » domande sull'una, sull'altra, o su ambedue le telescriventi, ed analizzando le risposte che riceve. La situazione è quindi, per il calcolatore, quella di una « prova » se gli si affida il compito di rendere vani gli sforzi di I: e cioè C deve rispondere alle domande di I in modo che I non possa distinguere C da M. Perciò la nuova domanda, che risulta dalla prova di Turing, e che sostituisce la primitiva: « La macchina può pensare? » è la seguente: « Vi sono macchine tali che si possa fare in modo che abbiano successo in questo gioco? ».

È chiaro che la domanda ha un carattere differente da quella primitiva, perché che cosa si intende per « aver successo » nel gioco? È una questione di grado, non materia di « sì » o « no ». Il criterio per giudicare del successo sarà probabilmente statistico: in media, l'interrogante I ha una probabilità del p per cento di identificare in modo corretto la macchina dopo avere interrogato per m minuti. Il fatto è che la prova di Turing pone problemi di progetto e di ingegneria, e in particolare di programmazione: e cioè: fino a che punto sappiamo fare una macchina che faccia questo gioco? che ci attendiamo di saper fare tra 10 o 20 anni? Si può sempre non essere d'accordo su tali questioni, ma l'emotività di cui era carica la domanda originale è quasi completamente scomparsa.

Nel suo scritto del 1950, Turing esprimeva la convinzione che, alla fine del secolo, si sarebbero potute costruire macchine capaci di fare così bene il gioco, che un interrogatore medio non avrebbe avuto più del 70% di probabilità di giungere alla corretta identificazione dopo cinque minuti di interrogazione.

Non ci risulta che alcuno abbia tentato veramente di fare in modo che una macchina possa affrontare la prova di Turing. Si è tuttavia compiuta una buona quantità di lavoro che getta luce su come si potrebbe fare a programmare un calcolatore per la prova di Turing, e che grado di successo ci si potrebbe atten-

dere, se un tentativo fosse fatto all'attuale stato dell'arte della programmazione euristica.

Poiché un calcolatore, per avere successo nella prova di Turing, deve poter conversare in modo sufficiente, per mezzo di parole stampate, in un linguaggio naturale, le più importanti indicazioni si ricavano da certi programmi costruiti per tenere conversazioni plausibili in ristretti campi di discorso. Discuteremo due di questi programmi scritti in anni recenti. Lasciamo che il lettore giudichi da sé, da questi esempi, quanto sia giustificata la previsione di Turing.

4. Programmi di conversazione

Vi sono calcolatori programmati per tenere conversazioni plausibili, e talvolta utili, in limitati domini di discorso. Una vivace attività di ricerca continua in tale area di programmi di calcolatori, che accettano e/o rispondono in linguaggio naturale.

Lo scopo di questo paragrafo è quello di dare al lettore una certa conoscenza del modo di operare di tali programmi, in modo che egli possa farsi una idea della relazione esistente fra ciò che un tale programma può compiere e la sua complessità. Egli potrà allora congetturare se e come si possa fare, con uno sforzo concentrato, ora o fra dieci anni, un programma che possa essere sottoposto alla prova di Turing.

Un programma che simuli un buon conversatore deve poter imitare due distinti modi di comportamento che si verificano nella conversazione umana. In primo luogo, esso deve capire una data affermazione o una data domanda in una qualunque delle forme in cui essa può essere formulata; deve, insomma, comprendere il significato di una frase. In secondo luogo, deve poter conversare in modo plausibile su un dato soggetto anche se non dispone dei dati necessari per dare una risposta chiara o per fare una affermazione. In breve, deve poter dare una risposta plausibile a una domanda come: « Che cosa pensa lei di Picasso? ».

La distinzione fra « capire » e « dare una risposta plausibile » è ben lontana dall'essere chiara e netta: noi cerchiamo di chiarirla con un'analogia fra due stati mentali di cui ogni lettore deve avere avuto esperienza. Da un lato, l'intenso sforzo per comprendere l'essenza di un argomento: dall'altro lo stato mentale di chi non è interessato in una discussione, ma che, per non sembrare scortese, accenna col capo e mormora frasi non compromettenti come: « Molto interessante! ».

È tipico dello stato dell'arte attuale il fatto che i programmi di conversazione finora scritti tendono ad accentuare l'uno o l'altro di questi due aspetti, ma non tutti e due. Siamo ancora nello stadio in cui i ricercatori tentano di mettere a punto una tecnica per volta, anziché introdurre diverse nello stesso programma.

Tra i numerosi programmi scritti negli ultimi anni, due, detti ELIZA e STUDENT, illustrano abbastanza bene i due estremi (rispettivamente « conversazione

con scarsa comprensione » e « comprensione con scarsa conversazione »).

ELIZA cerca di tenere una conversazione plausibile senza tentare di « capire » ciò di cui si parla. La parola « capire » è usata in questo senso: se un uomo si comportasse come è descritto qui sotto, noi diremmo certamente che egli rispondeva senza capire l'argomento discusso.

ELIZA è, in sé, un programma di carattere generale che accetta frasi in inglese e risponde con proprie frasi, ottenute dalle frasi in ingresso mediante certe trasformazioni. Queste trasformazioni sono determinate da un « copione », ed ELIZA si comporta in modo assai differente a seconda di quale « copione » sta seguendo. Secondo uno di tali copioni, detto DOCTOR, ELIZA risponde in modo molto simile a quello di uno psichiatra. Esso ha prodotto molte notevoli conversazioni, di cui la seguente è un esempio. (Le risposte di ELIZA sono in maiuscolo)

* * *

Gli uomini sono tutti uguali.

IN CHE MODO.

Ci stanno sempre seccando per una cosa o per un'altra.

PUÒ PENSARE A UN ESEMPIO SPECIFICO.

Bene, il mio ragazzo mi ha fatto venire qui.

IL SUO RAGAZZO LA HA FATTA VENIRE QUI.

Dice che sono sempre depressa.

MI RINCRESCE SENTIRE CHE LEI È DEPRESSA.

È vero. Io sono infelice.

LEI PENSA CHE IL VENIRE QUI L'AIUTERÀ A NON SENTIRSI INFELICE.

Io ho bisogno d'aiuto, questo è certo.

CHE COSA VORREBBE DIRE PER LEI AVERE AIUTO.

Forse potrei imparare ad andare d'accordo con mia madre.

MI DICA QUALCOSA DI PIÙ DELLA SUA FAMIGLIA.

Mia madre si prende cura di me.

CHI ALTRI NELLA SUA FAMIGLIA SI PRENDE CURA DI LEI.

Mio padre.

SUO PADRE.

Lei assomiglia in qualche modo a mio padre.

CHE RASSOMIGLIANZA VEDE.

Lei non è molto aggressivo ma penso che non voglia che io me ne accorga.

CHE COSA LE FA PENSARE CHE IO NON SIA MOLTO AGGRESSIVO.

Lei non discute con me.

PERCHÉ LEI PENSA CHE IO NON DISCUTA CON LEI.

Lei ha paura di me.

LE FA PIACERE CREDERE CHE IO HO PAURA DI LEI.

Mio padre ha paura di tutti.

COS'ALTRO LE VIENE IN MENTE QUANDO PENSA A SUO PADRE.

Sgridate.

HA QUALCOSA A CHE FARE QUESTO CON IL FATTO CHE IL SUO RAGAZZO LA HA FATTA VENIRE QUI.

* * *

Il metodo seguito da ELIZA per produrre risposte come quelle sopra riportate è descritto dal suo progettista, J. Weizenbaum del MIT, come segue:

« La procedura di massima del programma è molto semplice: il testo è letto ed esaminato per trovare una parola chiave. Se questa è trovata, la frase viene trasformata seguendo una regola associata con la parola chiave; se non, si richiama una osservazione senza contenuto, o, in certe condizioni, una trasformazione precedente. Il testo elaborato o richiamato viene poi stampato ».

Si prenda in esame la frase « Io sono molto infelice in questi giorni ». Supponiamo che uno straniero, con una limitata conoscenza dell'inglese, ma con un buon orecchio abbia sentito pronunciare questa frase ma abbia capito solo le prime due parole « Io sono ». Desiderando sembrare interessato, e forse anche comprensivo, egli può rispondere « Da quanto tempo Lei è stato infelice in questi giorni? ». Ciò che egli ha dovuto fare è stato di applicare alla frase originale una specie di maschera, una parte della quale coincide con le due parole « Io sono », e l'altra isola le parole « molto infelice in questi giorni ». Egli deve anche avere un « attrezzo di rimontaggio » associato specificamente a questa maschera, che prescrive che ogni frase della forma « Io sono BLAH » può essere trasformata in: « Da quanto tempo Lei è BLAH », qualunque sia il significato di BLAH. Un esempio un poco più complicato è dato dalla frase « Sembra che tu mi odii ». Ora lo straniero capisce solo le parole « tu » e « mi ». Egli applica la maschera che suddivide la frase nelle quattro parti:

1) Sembra che 2) tu 3) mi 4) odii

di cui egli capisce solo la seconda e la terza parte. La regola di rimontaggio potrebbe allora fornire:

« Che cosa ti fa pensare che io ti odii »

e cioè, può scartare la prima componente, tradurre le due parole note « tu » e « mi » in « io » e « ti », e attaccare una frase di uso generale « Che cosa ti fa pensare » davanti alla frase ricostruita.

Il « copione » DOCTOR consiste di circa una pagina e mezza di regole di decomposizione e ricomposizione come descritte sopra. È veramente sorprendente che un insieme di regole così relativamente semplice possa produrre una conversazione così credibile come quella riportata. Tuttavia, questa semplicità ha un prezzo. Nel copione discusso, uno degli scopi di ELIZA è quello di nascondere la sua mancanza di comprensione. Se nessuna regola di decomposizione può essere applicata, la risposta può essere una frase come « Parliamo di qualcos'altro. Dimmi qualcosa di questo e di quello ». Una frase come questa, che non richiede alcuna comprensione della discussione precedente, può sempre entrare nella conversazione, purché non troppo spesso.

Il secondo programma in discussione, detto STU-

DENT, risolve problemi verbali di algebra come i seguenti:

« La somma di tre numeri è 9. Il secondo numero è il doppio del primo numero, più tre. Il terzo numero è uguale alla somma dei primi due numeri. Trovare i tre numeri ».

« Maria ha l'età doppia di quella che aveva Anna quando Maria aveva l'età che ha Anna ora. Se Maria ha 24 anni, quanti anni ha Anna? ».

Quando gli viene presentato il primo problema, STUDENT stampa:

IL PRIMO NUMERO È 0,5

IL SECONDO NUMERO È 4

IL TERZO NUMERO È 4,5.

Per il secondo problema, esso stampa:

ANNA HA 18 ANNI.

Per quanto riguarda STUDENT, il mondo consiste di descrizioni verbali di equazioni algebriche, e, in particolare, di equazioni di primo grado. Quando una descrizione di questo genere gli viene presentata, STUDENT procede ad identificare le incognite contenute e le relazioni fra esse, e a tradurre la descrizione verbale in un insieme di equazioni nella notazione matematica normale.

Nel primo problema, STUDENT identifica le tre incognite p (primo numero), s (secondo numero) e t (terzo numero) e poi stabilisce il sistema di equazioni:

$$p + s + t = 9$$

$$s = 2p + 3$$

$$t = p + s$$

Giunto a questo punto, il resto è ovvio, dato che vi sono modi ben noti per risolvere sistemi di equazioni lineari. Quello che ci interessa in STUDENT sta tutto nella sua capacità di tradurre frasi da un ristretto, ma informale, sottoinsieme della lingua inglese, in equazioni algebriche espresse in notazione formale. La nostra discussione di come STUDENT può eseguire tale traduzione è necessariamente semplificata, ma speriamo di rendere chiare le idee e le tecniche principali.

STUDENT suddivide le parole e le espressioni nelle frasi che analizza in tre categorie: variabili, operatori e sostituenti.

I sostituenti sono espressioni che STUDENT sostituisce immediatamente con altre. Per esempio, la parola « doppio » è sempre sostituita da « 2 volte ». Lo scopo di tali sostituzioni è quello di ridurre il numero delle forme linguistiche, col sostituire molte diverse e speciali costruzioni con un minor numero di costruzioni più generali.

Gli operatori sono espressioni che denotano le operazioni aritmetiche che STUDENT può compiere (e cioè addizione, sottrazione, moltiplicazione, divisione, elevazione a potenza) o combinazioni di tali operazioni (come la somma di x col prodotto di y e z), e, infine, la relazione di uguaglianza.

Qualcuna delle forme linguistiche che possono essere usate per esprimere le operazioni aritmetiche e la relazione di uguaglianza sono elencate nella tabella:

Sommario delle forme linguistiche per esprimere le funzioni aritmetiche e la relazione di uguaglianza

$x=y$ x è y ; x uguaglia y ; x è uguale a y .

$x+y$ x più y ; la somma di x ed y ; x aggiunto ad y .

$x-y$ x meno y ; la differenza fra x ed y ; x tolto da y .

$x \cdot y$ x volte y ; x moltiplicato y ; $x y$ (se x è un numero).

x / y x diviso per y ; x per ciascun y .

La terza categoria di parole ed espressioni identificata da STUDENT è quella delle variabili. Una variabile è una sequenza di parole che STUDENT prende come rappresentante di una incognita del problema presentatogli. STUDENT non cerca affatto di comprendere il significato di una variabile: egli non fa altro che identificare quali espressioni sono variabili, e quali frasi indicano la stessa incognita. Se STUDENT può risolvere un problema verbale in cui si presenta la variabile « l'età di Anna », esso può anche risolvere un problema simile in cui ogni espressione « l'età di Anna » è sostituita da una parola senza senso come « fgh ». Per il fatto che STUDENT non cerca di comprendere alcunché circa la natura di una data variabile, risulta ovvio che la capacità di STUDENT di riconoscere differenti espressioni che indicano la stessa variabile è assai limitata. Per esempio, STUDENT non riconoscerebbe che « il numero di anni trascorsi dalla nascita di Anna » indica la stessa incognita che « l'età di Anna ».

Descriviamo ora il metodo principale usato da STUDENT per riconoscere le equazioni algebriche in travestimento verbale. Dopo aver sostituito tutti i sostituenti, STUDENT ricerca se vi sono operatori, confrontando le espressioni delle frasi di ingresso con le espressioni di operatore registrate in un catalogo. Poi identifica come variabili le sequenze di parole comprese fra gli operatori e decide quali sequenze indicano la stessa incognita. (Come si è detto prima, STUDENT riconosce due espressioni come indicanti la stessa incognita solo se esse sono identiche o differiscono di molto poco).

Se le frasi di ingresso sono del tipo che STUDENT può comprendere, a questo punto STUDENT ha costruito un sistema di equazioni algebriche in notazione matematica formale. Se sa risolvere queste equazioni, esso procede alla loro soluzione e stampa i risultati. In caso contrario, per esempio se qualche equazione non è lineare, o se vi sono meno equazioni che incognite, STUDENT domanda: « Conoscete altre relazioni fra queste variabili? ». Se la risposta è « No », esso stampa: « Non so risolvere questo problema ».

È notevole il fatto che questo modo di affrontare il problema permette a STUDENT di fare con successo quello che fa. Naturalmente, se si è capito il modo di operare di STUDENT è facile costruire esempi che lo fanno sbagliare. Per esempio, in un problema che contenga la variabile « il numero di volte che le vidi », STUDENT identifica subito

« volte » come l'operatore di moltiplicazione e ne deduce che vi sono due variabili: « il numero di » e « che la vidi ».

Tuttavia, i problemi che STUDENT ha risolto sembrano indicare che, entro il suo ristretto campo algebrico, esso può competere con lo studente medio delle scuole medie superiori. Ciò dà un certo peso alla conclusione di D.G. Bobrow, che ha scritto STUDENT: « Io penso che siamo ancora lontani dal saper scrivere un programma che possa capire tutta, o almeno in gran parte, la lingua inglese. Tuttavia, nel suo ristretto campo di competenza, STUDENT ha dimostrato che possono essere costruite macchine che "capiscono" ».

5. Può una macchina riprodurre se stessa?

L'autoriproduzione, e cioè la capacità di un organismo di produrne altri simili a sé, è comunemente considerata una prerogativa degli organismi viventi, e un carattere indicativo della distinzione fra materia vivente e materia non vivente. I recenti risultati della sintesi in laboratorio di molecole complesse capaci di riprodursi sono stati indicati come creazione di « vita in provetta », attestando ancora una volta che, comunemente, vita e autoriproduzione sono concetti inseparabili. Tuttavia, il fatto che vi sono, in natura, diversi fenomeni molto simili all'autoriproduzione e relativi a esseri non viventi, (per esempio, l'accrescimento di certi cristalli), suggerisce che la distinzione fra cose che si riproducono e cose che non lo fanno non sia necessariamente la stessa che fra cose viventi e cose non viventi.

Un argomento apparentemente plausibile, secondo cui appare impossibile costruire una macchina autoriproducibile, è il seguente: se una macchina A costruisce una macchina B, essa deve certamente contenere una descrizione completa di B; ed inoltre A deve comprendere certe altre parti, quali un dispositivo di comando capace di interpretare tale descrizione, organi meccanici capaci di effettuarne la costruzione, ecc. Perciò è chiaro che A deve essere più complessa di B, e che nessuna macchina può costruirne un'altra della stessa complessità.

L'esperienza delle macchine che conosciamo sembra confermare questo argomento. Per esempio, in una linea di montaggio, una macchina che non faccia altro che avvitare a fondo delle viti deve essere sostanzialmente più complessa di una vite. Benché appaia intuitivamente convincente, tale argomento è ingannevole, come mostrerà la discussione seguente. Noi abbiamo molta esperienza di macchine relativamente semplici, ma la nostra intuizione di che cosa possano fare, in via di principio, le macchine, e soprattutto macchine molto complesse, non è granché buona. John von Neumann si è occupato di tale materia e l'ha studiata a fondo negli anni '50. Egli ha suggerito che la capacità di un organismo di riprodurre se stesso sia una indicazione della complessità di tale organismo, e si è occupato dello studio degli « automi autoriproducibili », che possono essere de-

finiti con rigore e completezza, piuttosto che degli organismi viventi, che sono così complessi che nessuno li conosce completamente.

Nel suo studio, von Neumann ha preso in considerazione alcuni differenti modelli di macchine, due dei quali saranno discussi qui: il primo è un dispositivo meccanico, costruito in modo abbastanza tradizionale, con un montaggio di parti elementari; il secondo è un modello algebrico astratto. Il primo modello ha il vantaggio di richiamarsi intuitivamente all'esperienza che abbiamo di dispositivi meccanici, e perciò la costruzione di una macchina autoriproducibile può essere facilmente capita nelle sue idee essenziali. Ha lo svantaggio che la sua descrizione dettagliata, per esempio sotto forma di disegni costruttivi, sarebbe assai complessa e richiederebbe un tempo eccessivo. Il secondo modello cerca di evitare di preoccuparsi eccessivamente dei particolari e presenta in modo semplificato l'idea dell'autoriproduzione. È allora possibile effettuare un progetto particolareggiato di una macchina autoriproducibile. Von Neumann ha fatto ciò in un manoscritto di circa 200 pagine. Lo svantaggio del secondo modello, agli scopi espositivi, è che esso richiede un notevole sforzo di pensiero per capire in quale senso precisamente questo modello algebrico contiene la nozione dell'autoriproduzione della macchina. Noi perciò dedicheremo solo poche parole a tale modello, e daremo poi un profilo della costruzione di una macchina autoriproducibile secondo il primo modello di von Neumann. Nel modello algebrico, usualmente detto il modello « cellulare » o « a mosaico », lo spazio è rappresentato come uno schieramento infinito bidimensionale e rettangolare di celle identiche, e si ammette che il tempo proceda per intervalli discreti. Ogni cella è essa stessa una semplice macchina astratta, che, ad ogni istante di tempo, può assumere uno stato fra un numero finito di stati, ed è capace di scambiare segnali con le quattro celle adiacenti. V'è una funzione di transizione che determina il comportamento nel tempo di ogni cella, prescrivendo un cambiamento di stato da un istante al successivo, che dipende solo dallo stato presente della cella e delle sue quattro vicine. Uno degli stati possibili di ogni cella è detto « stato ineccitabile »: una cella in tale stato può considerarsi una cella vuota. Nel modello algebrico una macchina si identifica con una configurazione geometrica finita di celle contigue e con la specificazione dello stato di ogni cella, ed opera come segue. Prima di un istante iniziale t_0 , tutte le celle dello schieramento infinito sono nello stato non eccitabile: al tempo t_0 tutte le celle della detta configurazione geometrica sono portate in stati specificati: le altre restano come sono. Da allora in poi, l'operazione è comandata, passo a passo, dalla funzione di transizione, che determina i cambiamenti individuali di stato per ogni cella. Si può allora chiamare autoriproducibile una macchina (una configurazione di celle contigue in dati stati), se ad un certo tempo finito dopo t_0 vi sono due configurazioni

finite contenute nello schieramento infinito, ognuna uguale alla macchina originale.

Descriviamo ora il primo modello di macchina autoriproducibile, il dispositivo meccanico. Si immagini un grande magazzino contenente una grande, potenzialmente illimitata, provvista di un numero finito di parti meccaniche. Esse sono le « parti elementari » di cui è composta la macchina, e debbono essere di una varietà tale da permettere di costruire con esse veicoli mobili e provvisti di sorgenti autonome di energia. Esse possono essere, per esempio, ruote, ingranaggi, alberi, viti e dadi, nastri di carta, batterie, ecc. Come vedremo, la descrizione precisa di questa provvista di parti può presentare difficoltà, ma, per il momento, supponiamo che una tale raccolta di parti possa essere convenientemente definita. La questione della possibile esistenza di una macchina autoriproducibile in questo ambiente può essere formulata come segue: dato un inventario sufficientemente ricco di parti elementari, può esistere una macchina, composta di tali parti, che possa aggirarsi per questo magazzino, prelevare le parti dagli scaffali, e montarle assieme in modo da costruire una copia di se stessa, e cioè una seconda macchina identica, che possa poi procedere a costruire una terza copia, e così via?

Ora dobbiamo occuparci della precisa descrizione della provvista delle parti. È abbastanza semplice evitare il problema mediante una improbabile scelta delle parti elementari. Per esempio, la questione diventa banale se una delle parti elementari è definita come l'intera macchina senza la batteria, o senza qualcosa di egualmente semplice: in tale caso l'autoriproduzione consiste semplicemente nella installazione della batteria. Si resta del tutto insoddisfatti da un argomento di questo tipo, perché esso assume per una parte elementare essenzialmente le stesse proprietà che si cerca di stabilire per la macchina completa. D'altra parte, se le parti elementari si scelgono troppo piccole o troppo semplici, allora le proprietà essenziali della macchina possono venire nascoste dai particolari della costruzione mediante tali parti. Non vi sono quindi regole rigorose, ma solo il senso comune, per suggerire la definizione delle parti elementari: il criterio per stabilire la opportunità di una certa scelta sembra essere quello di stabilire se il problema che ne risulta è o no « interessante ». Lo stesso von Neumann ha descritto, una volta, un gruppo di una dozzina di parti elementari che egli pensava potessero essere usate per costruire la macchina, benché non ne abbia effettuato la descrizione in termini di tali parti. Egli ha invece suddiviso la macchina autoriproducibile in tre macchine più piccole, nessuna delle quali è capace di riprodurre se stessa, e che sono tutte macchine plausibili. Queste tre macchine sono: 1) un costruttore universale U; 2) un duplicatore di nastro D; 3) un automa di governo C. Non descriveremo la laboriosa costruzione di tali macchine da una raccolta di parti elementari, ma facciamo appello alla intuizione del lettore che tale costruzione è possibile in via di princi-

pio. Descriveremo invece senz'altro queste macchine e il loro funzionamento.

Il costruttore universale U è la parte più critica del modello di von Neumann di una macchina autoriproducibile. U è una macchina che, data una descrizione $\mathcal{D}$ (M) di una macchina qualunque, composta delle parti $P_1...P_k$, può fabbricare con tali parti un esemplare della macchina M. L'analogia con un calcolatore universale (*general purpose*) che, data una sufficiente capacità di memoria può essere programmato per calcolare ogni funzione « calcolabile », convincerà la maggior parte di coloro che hanno familiarità con i calcolatori che una tale macchina può essere costruita. Un impianto industriale comandato da un calcolatore e provvisto di linee di montaggio è abbastanza simile alla macchina U. Un progetto completo di U, con la specifica della lista delle parti, della descrizione, di un algoritmo di costruzione, appare certamente come un notevole sforzo di ingegneria, ma rientra del tutto nelle capacità della tecnologia moderna.

Il costruttore universale U richiede, per poter funzionare una descrizione non ambigua, in forma leggibile da macchina, della macchina che costruisce. Ciò solleva un altro problema nella scelta della classe delle parti elementari: e cioè, le parti debbono essere scelte in modo da essere sicuri che esiste una descrizione finita di ogni macchina M che può essere costruita con esse.

Ciò non impone in pratica una grande restrizione, poiché è perfettamente ragionevole supporre un insieme di parti tali che, se due parti debbono essere unite in qualche modo, esse possano essere unite solo in un numero finito di modi.

Una volta scelta una conveniente classe di parti, non vi sono speciali difficoltà nel progettare un linguaggio formale che descriva le parti stesse e il modo in cui debbono essere interconnesse. Supporremo che una descrizione $\mathcal{D}$ (M) di una macchina M sia codificata in forma di una banda di carta perforata, ma qualunque descrizione leggibile da macchina è soddisfacente.

Le altre due macchine componenti sono assai più semplici di U. Il duplicatore di zona D è semplicemente una macchina che legge una banda di carta perforata e ne fa una copia perforando l'identica configurazione di fori su una banda vergine. L'automa di governo, C, è una macchina semplice, che può comandare il funzionamento di U e di D, e cioè avviarle e fermarle, inserirvi zone, ecc. La funzione specifica di C è chiarita dalla descrizione del comportamento della macchina composta $U+D+C$.

Sia M una macchina qualunque, e $\mathcal{D}$ (M) la sua descrizione in termini di configurazione $P_1...P_k$, perforate sulla banda di carta. $\mathcal{D}$ (M) sia unita alla macchina $U+D+C$. Dapprima la macchina di governo C inserisce la banda $\mathcal{D}$ (M) nel duplicatore D, ottenendo due copie di $\mathcal{D}$ (M). Poi C inserisce una copia di $\mathcal{D}$ (M) nel costruttore universale U, che quindi costruisce la macchina M secondo la descrizione. Infine, C fornisce l'altra copia

di $\mathcal{D}$ (M) alla macchina M appena fabbricata, e riunisce la descrizione originale $\mathcal{D}$ (M) a $(U+D+C)$. Così, per una macchina qualunque M, $U+D+C+\mathcal{D}$ (M) costruisce $M+\mathcal{D}$ (M). Per l'autori-produzione, quindi, tutto ciò che si richiede è di fornire alla macchina $U+D+C$ la banda $\mathcal{D}$ ($U+D+C$): allora $U+D+C+\mathcal{D}$ ($U+D+C$) costruisce precisamente una copia di se stessa.

Qui termina la nostra discussione dei modelli di macchine autoriproducenti di von Neumann. Abbiamo visto una macchina operante in un ambiente interessante, cioè un ambiente in cui le parti elementari possono essere semplici e familiari come quelle che incontriamo nella vita di ogni giorno, eppure tale macchina ha la capacità non comune di autoriprodursi, il che, nella nostra esperienza, abbiamo visto accadere solo negli esseri viventi. Quindi alla domanda posta all'inizio di questo paragrafo si può rispondere affermativamente.

6. Note e citazioni bibliografiche

L'articolo qui presentato costituisce parte della memoria: G. NIEVERGELT, J.C. FARRAR, *What machines can and cannot do*, Computing Surveys, Vol. 4, No. 2, 1972.

Il falso automa scacchista di Von Kempelen è citato spesso nella letteratura. Si dice che si debba in parte ad esso lo sviluppo di interesse a tale gioco in America nella prima parte del diciannovesimo secolo. La sua interessante storia è sommariamente tracciata in:

R.K. HAGEDORN, *Benjamin Franklin and chess in early America*, Univ. Pennsylvania Press, Philadelphia Pa, 1958.

L'interesse per i programmi di gioco degli scacchi è stato vivo da quando esistono i calcolatori elettronici. Lo sviluppo di tali programmi costituisce un'altra illustrazione del come i calcolatori vengano programmati per eseguire, in modo non banale, una attività che è generalmente considerata di alto livello intellettuale. La prima comunicazione che descrive come possa essere scritto un programma di scacchi è:

C.E. SHANNON, *Programming a digital computer for playing chess*, Philosophy Magazine 41 (March 1950), 356-357.

In questo articolo, Shannon accenna anche alla questione dell'intelligenza delle macchine, ed esprime idee simili a quelle di Turing. Questo e diversi altri articoli degni di nota sono riprodotti nella raccolta:

J.R. NEWMAN (Ed.), *The world of mathematics*, vol. 4, Simon and Schuster, New York, 1956.

La descrizione di uno dei recenti programmi di scacchi detto MAC HACK VI, e di alcuni semplici giochi si trova in:

R.D. GREENBLATT, D.E. EASTLAKE, S.D. CROCKER, *The Greenblatt chess program*, in Proc. AFIPS 1967 FJCC vol. 31 AFIPS Press, Montvale, N.J. 801-810.

Fra i programmi di scacchi finora scritti il maggior successo l'ha avuto il programma di Samuel, che ha raggiunto il livello di « Maestro ». E' descritto nei due articoli seguenti:

A. SAMUEL, *Some studies in machine learning using the game of checkers*, IBM J. R & D 3, 3 (July 1954) 210-229; e

A. SAMUEL, *Some studies in machine learning using the game of checkers. II Recent progress*, IBM J. R & D. 11,6 (Nov. 1967) 601-617.

La discussione originale di Turing sulla questione se una macchina può pensare è apparsa in

A.M. TURING, *Computing machinery and intelligence*, Mind 9 (1950) 433-460

ed è ristampato in *The world of mathematics*, citato sopra (pp. 2099-2123).

Lo scritto di Turing, che è generalmente considerato un classico della letteratura, fu uno dei primi contributi ad alto livello ad una vivace discussione del problema mente-macchina. Per una buona antologia degli scritti filosofici sul soggetto, che comprende lo scritto di Turing e anche qualche critica ad esso, ed è utile per porre nella giusta prospettiva il problema del pensiero delle macchine, vedi:

R. ANDERSON (Ed.), *Minds and machines*, Prentice Hall Inc. Englewood Cliffs, N.J. 1964.

Una descrizione di diversi importanti programmi atti a rispondere a domande poste in linguaggio naturale si trova in

M. MINSKY (Ed.), *Semantic information processing*, MIT Press, Cambridge Mass. 1968.

Il capitolo 3 di tale libro descrive il programma risolutore di problemi algebrici STUDENT, che fu sviluppato in una tesi di laurea al MIT:

D.G. BORROW, *Nature language input for a computer problem-solving system in Semantic information processing*, 135-215.

Il programma di conversazione ELIZA è stato descritto per la prima volta in

J. WEIZENBAUM, *ELIZA - a computer program for the study of natural language communication between man and machine*, Comm. ACM 9, 1 (Jan. 1966) 36-45 e più tardi in

J. WEIZENBAUM, *Contextual understanding by computers*, Comm. ACM 10, 8 (Aug. 1967) 474-480.

La discussione di John von Neumann sulle macchine autoriproducenti è parte di uno scritto che espone i suoi concetti su una teoria matematica degli automi. E' apparso sotto il titolo *The general and logical theory of automata* in

L.A. JEFFRESS, *Cerebral mechanisms in behavior - the Hixon symposium*, John Wiley & Sons, New York 1951, 1-41. E' anche ristampato nel sopracitato *The world of mathematics* (pp. 2070-2098) e in

A.H. TAUB (Ed.), *John von Neumann - collected works*, vol. V, Macmillan Co., New York, 1963, 288-328.

Il modello a mosaico di autoriproduzione menzionato nel testo è descritto particolareggiatamente da von Neumann in un manoscritto che fu incominciato nel 1952 e lasciato incompleto al momento della sua morte. Questo manoscritto fu successivamente completato e curato da A.W. Burks, e pubblicato come parte di

J. VON NEUMANN, *Theory of self-reproducing automata (Edited and completed by Arthur W. Burks)*, Univ. Illinois Press, Urbana, Ill., 1966.

L'attività di ricerca in tale campo continua, e un sommario recente si trova in

A.W. BURKS (Ed.), *Essays on cellular automata*, Univ. Illinois Press, Urbana, Ill., 1970.

Rilievi sul mercato dei calcolatori in Italia

RENATO LEVRERO

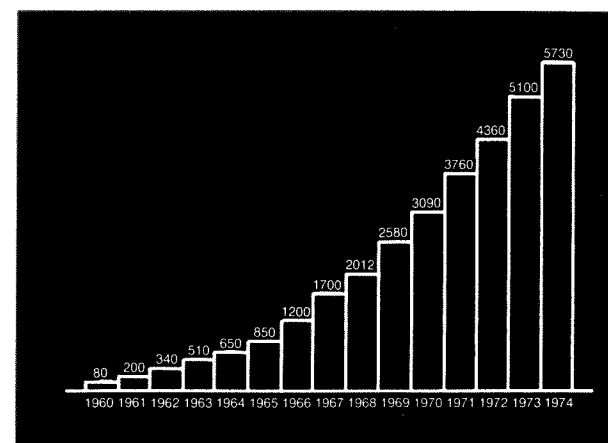
Honeywell Information Systems Italia
Servizio Studi Economici e Ricerche di Mercato

1. Introduzione

Quale sia la situazione delle conoscenze sulla struttura del mercato dei calcolatori elettronici in Italia, è noto da un precedente articolo [1] pubblicato su questa stessa rivista. La mancanza di rilevazioni ufficiali od ufficiose da parte di Enti Pubblici sul numero e valore del parco dei calcolatori in Italia è andata, a ben vedere, aumentando col passaggio dagli anni '60 agli anni '70. È infatti del 1968 la pubblicazione, da parte dell'ISTAT [2], dell'unico studio ufficiale sul parco installato in Italia (riferito al 31 marzo di quell'anno), mentre è del 1969 lo studio della Cassa di Risparmio delle Province Lombarde [3] che illustra la dinamica delle installazioni in Italia ed in Lombardia a partire dal 1955.

Dopo queste iniziative sarebbe parso logico attendersi un rinnovato impegno da parte degli Enti Pubblici di Ricerca e da parte dell'Amministrazione per la conoscenza dell'evoluzione di questo settore, in sintonia con quanto diversi Istituti specializzati andavano da tempo facendo negli altri paesi industrializzati; ma, al contrario, entrando negli anni '70,

Fig. 1 - Numero del parco installato di calcolatori general purpose in Italia. (Fonte: HISItalia)



le informazioni provenienti dal settore pubblico si sono andate via via rarefacendosi per poi sparire, negli ultimi anni, quasi del tutto. Le conoscenze sulla evoluzione di questo settore sono andate dunque progressivamente degradandosi: la stessa AICA, che nel 1971 [4] pubblicò una stima della struttura e consistenza del parco elaboratori installato in Italia, ha cessato di fornire l'aggiornamento sull'ulteriore evoluzione di quest'ultimo.

Le informazioni sono mancate, quindi, proprio nel momento in cui le oscillazioni della congiuntura rendono indispensabile una maggiore e più puntuale conoscenza della domanda ai fini di una razionale ed efficiente allocazione delle risorse, vuoi da parte degli operatori economici (pubblici e privati), vuoi da parte degli utenti.

Seppure finalizzato ad obiettivi di pianificazione interna, il Servizio Ricerche di Mercato della HISItalia, compie annualmente studi settoriali e rilevazioni sull'andamento del parco italiano che possono risultare di interesse generale ai fini di una conoscenza più puntuale della struttura della domanda in questo settore. Il fine che ci proponiamo in questo articolo, è quello di presentare una serie di informazioni statistiche sulla consistenza e struttura del parco elaboratori elettronici in Italia, cercando di mettere in relazione la dimensione e struttura del parco con quella dei principali aggregati economici, nonché tentando, ove possibile, qualche confronto con l'esperienza di paesi a più elevato indice di utilizzazione di materiale EDP. La grandezza alla quale faremo riferimento, per valutare il volume del parco, sarà in genere quella di parco « potenziale », dato dal numero degli elaboratori installati più gli ordini netti, in quanto questo concetto esprime meglio di quello di parco installato il volume della domanda effettiva di elaboratori presente sul mercato.

Nel grafico 1 diamo l'andamento storico del parco dei calcolatori installati in Italia a partire dal 1960, in numero di sistemi; il tasso medio di sviluppo del

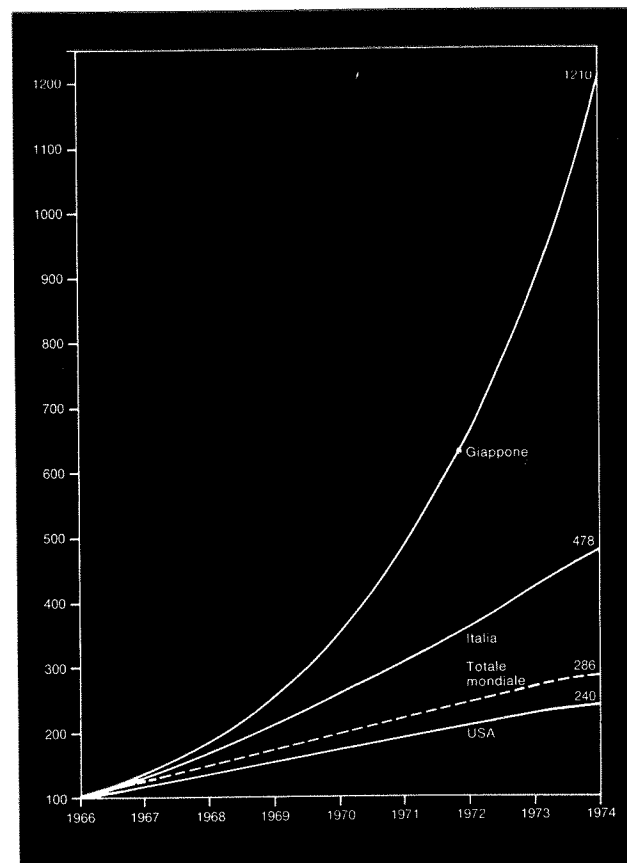


Fig. 2 - Sviluppo del parco mondiale di elaboratori e di quello italiano (1966 = 100). Fonte: HISItalia e per il Giappone JIP DEC REPORT N. 22 (1974).

parco è stato, nel primo quinquennio (1960-1965) del 43,9% medio annuo, mentre nel secondo quinquennio è stato del 25,9% annuo ed infine, per gli anni dal 1971 al 1974 del 15% circa annuo). Due considerazioni si impongono. In primo luogo si nota come il tasso di aumento del numero di sistemi installati sia assai più elevato di quello del prodotto nazionale lordo (per il periodo complessivo è stato attorno al 4,5%) ed è praticamente indipendente dall'andamento congiunturale di quest'ultimo. Secondariamente (vedi grafico 2), lo sviluppo del parco italiano risulta notevolmente più sostenuto della progressione del parco mondiale, dati soprattutto i più bassi livelli di saturazione che presenta il mercato italiano.

Complessivamente, dunque, nel 1974 erano installati in Italia 5730 calcolatori, con un aumento del 12,5% rispetto al 1973: aumenti dello stesso ordine si sono avuti in Francia e Germania, mentre più contenuti sono stati quelli dell'Inghilterra e degli USA. Mentre è possibile indicare con precisione il numero dei sistemi installati, per quanto riguarda il loro valore la cifra è più approssimata: nel 1974 un valore di locazione annua attorno ai 450 miliardi di lire dovrebbe costituire una buona approssimazione.

Sotto questo punto di vista, l'aumento del valore del parco rispetto al 1973 è stato del 16,6% confrontato con l'1,1% degli USA, il 12% del Canada e del-

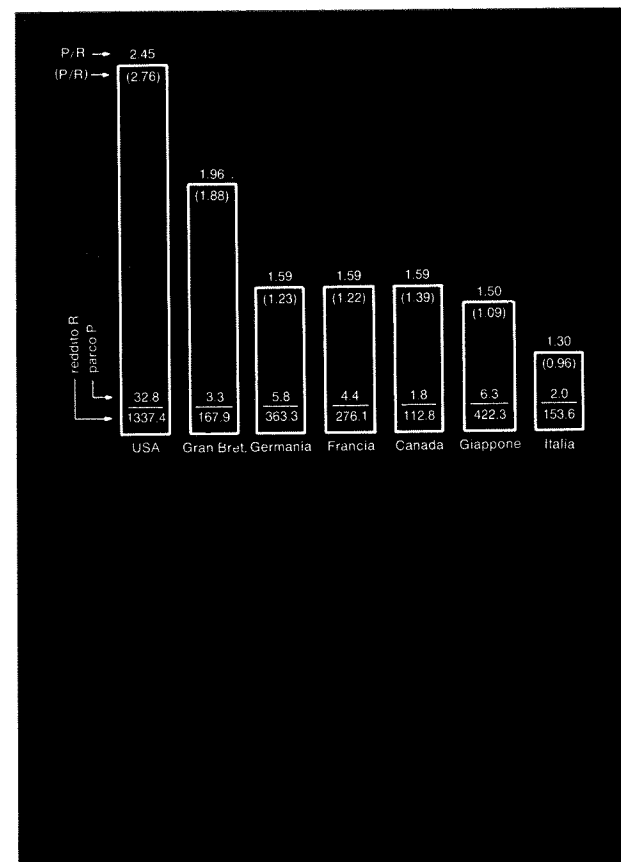


Fig. 3 - Rapporto ($\times 100$) tra il valore P (« if sold value ») del parco calcolatori installato ed il reddito interno R dei settori non agricoli (P ed R in miliardi di \$). I dati si riferiscono al 1974. Tra parentesi invece il rapporto (P/R) relativo al 1971. (Fonte: OCDE per il reddito, HISItalia per il parco per il 1974, IDC per il 1971).

L'Inghilterra, il 14% della Germania e della Francia ed il 21% del Giappone.

Per valutare il peso relativo che hanno nei diversi paesi industrializzati i calcolatori sull'insieme dell'economia, proponiamo due indicatori, abbastanza differenti da quelli tradizionalmente in uso: il primo (vedi fig. 3) è il rapporto fra la spesa per l'hardware ed il prodotto nazionale dell'industria, Servizi e Pub-

Fig. 4 - Occupazione dipendente extragricola e numero di elaboratori installati (1974). (Fonte: OCDE per gli occupati e HISItalia per il parco).

Descrizione	Occupati (000)	Calcolatori (n.)	Calcolatori per ogni 10.000 occupati
USA	72.800	65.000	8.9
Francia	16.000	9.300	5.8
Germania	21.500	11.500	5.3
Italia	11.600	5.700	4.9
Regno Unito	21.600	8.400	3.9
Giappone	34.000	23.500	6.9
Svizzera	2.700	2.200	8.2

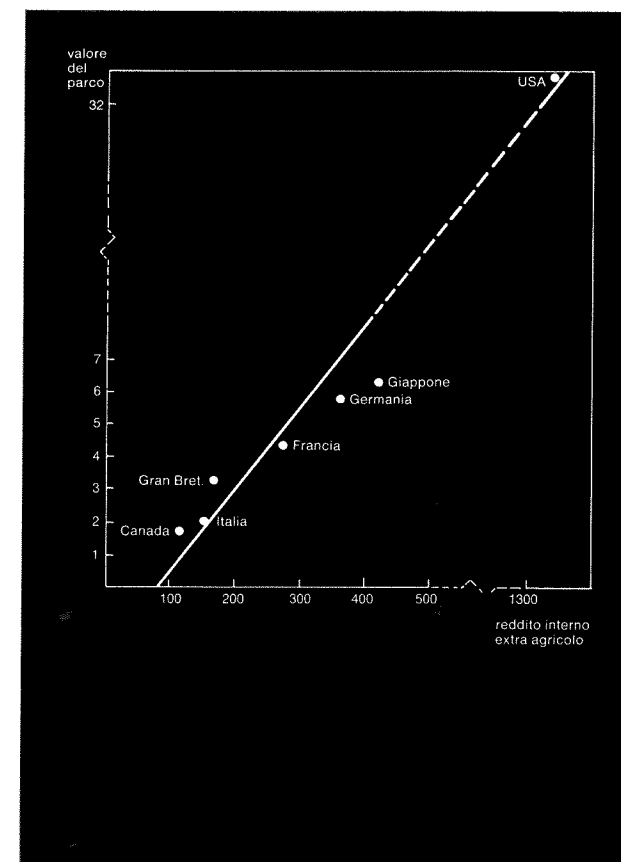


Fig. 5 - Correlazione tra valore del parco calcolatori e reddito interno extragricolo (in miliardi di \$).

blica Amministrazione, mentre il secondo (fig. 4) è il rapporto fra numero di calcolatori e l'occupazione dipendente nei settori extra-agricoli. Resta confermata, da questi indicatori, la parziale arretratezza dell'Italia nell'uso delle tecniche di elaborazione automatica dei dati rispetto agli altri paesi industrializzati: la spesa per hardware, infatti, in Italia è l'1,3% del reddito dei settori extra-agricoli, contro l'1,6% della maggior parte degli altri paesi (con l'eccezione dell'Australia). Si noti ancora come l'elevatissimo coefficiente di determinazione ($R^2 = 99\%$)⁽¹⁾ della regressione fra la spesa per materiale EDP ed il prodotto extra-agricolo confermi lo stretto rapporto esi-

Fig. 6 - Struttura della spesa per informatica (1974) in alcuni paesi. (Fonte: HISItalia).

Descrizione	Giappone	Australia	U.K.	Italia	Germania	Francia	Canada	USA
Spesa per:								
Unità centrali	39.4	30.6	36.0	37.0	30.8	27.3	39.4	35.6
Terminali	5.4	0.3	2.2	1.1	3.3	3.1	4.4	6.9
Data Entry		0.3	4.4	2.7	3.5	3.3	2.9	2.2
Software in house		22.8	19.9	22.9	21.2	19.2	15.8	17.2
Servizi software	32.4	10.2	2.8	5.0	5.3	7.3	1.6	0.8
Programmi		2.3	1.7	2.5	7.8	6.7	0.9	0.8
Servizi esterni		14.5	11.4	8.7			12.1	13.8
Costo operativo di eserc.	18.5	14.3	13.0	11.7	22.2	27.2	16.4	18.3
Consumi (schede, energ. ecc.)	4.3	4.7	7.8	8.4	5.9	5.9	6.5	4.4
Totale	100.0	100.0	100.0	100.0	100.0	100.0	100.0	100.0

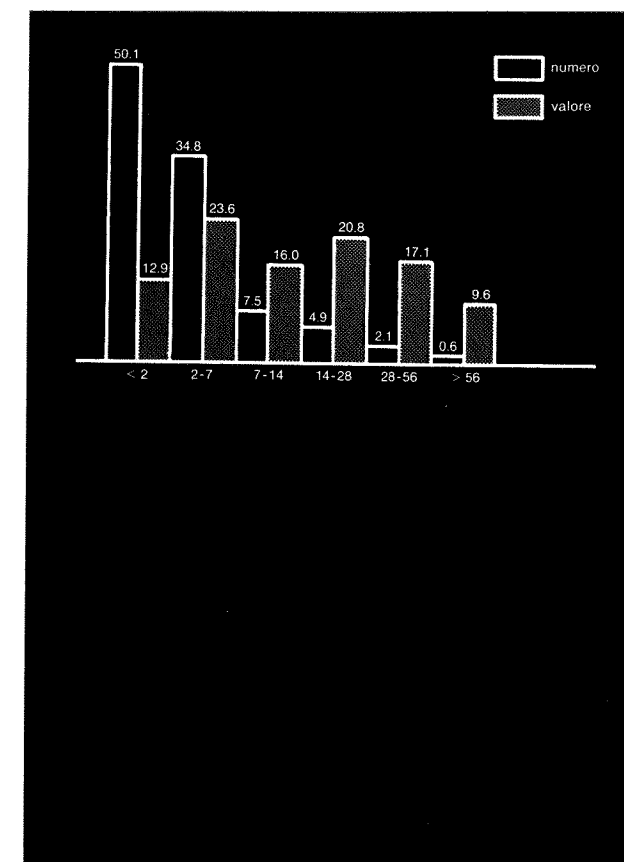


Fig. 7 - Composizione (%) del parco installato di calcolatori in Italia per classe di locazione (1974). (Fonte: HISItalia).

stente fra sviluppo industriale e meccanizzazione EDP (vedi grafico 5).

Prima di passare ad illustrare la struttura del parco italiano per settori ed aree geografiche, vogliamo concludere il paragone con altri paesi, esaminando la struttura della spesa complessiva per la mecca-

(1) L'indice di determinazione R^2 misura il grado di dipendenza lineare tra la variabile dipendente e la (o le) variabile (i) indipendente. E' pari al quadrato del coefficiente di correlazione e corrisponde alla quota di varianza spiegata dalla dipendenza ipotizzata. Nel nostro caso, ad esempio, le variazioni del prodotto interno extragricolo spiegano pressoché interamente (al 99%) le variazioni nell'ammontare del valore del parco installato nei vari paesi.

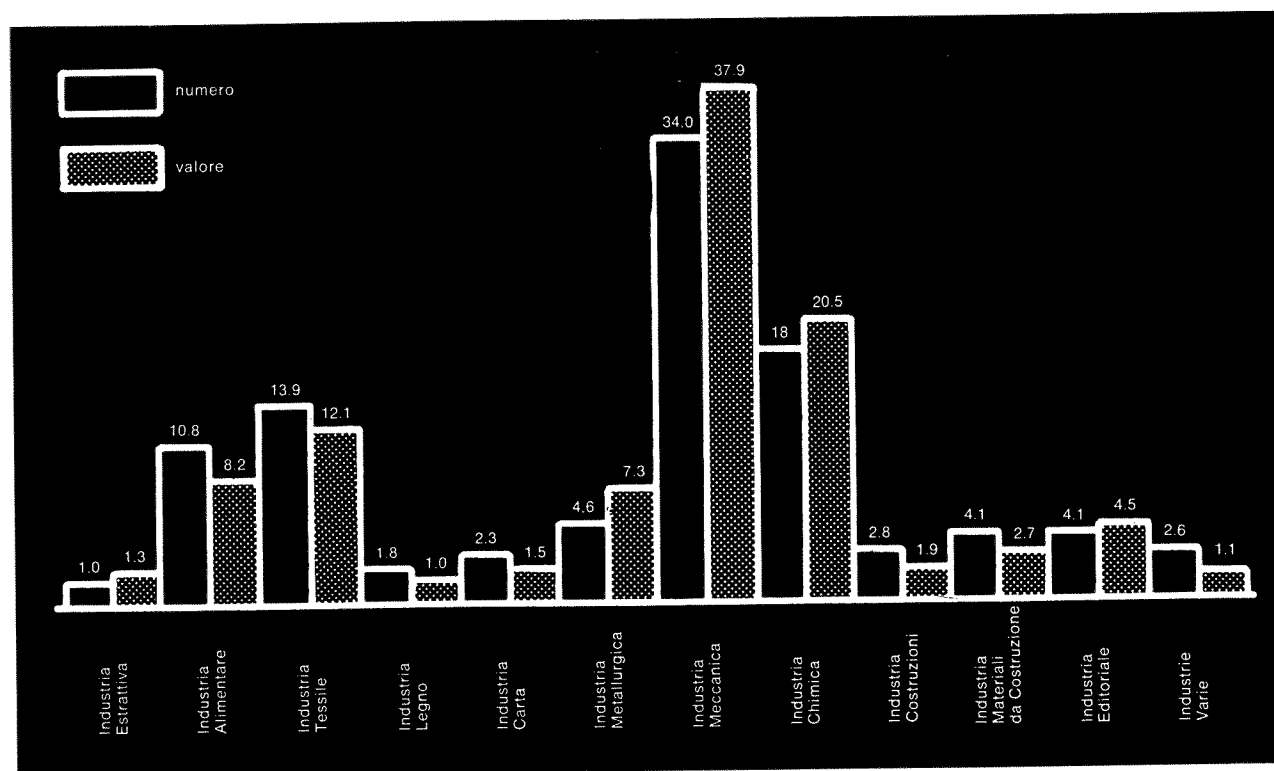


Fig. 8a) - Composizione (%) del parco calcolatori italiano per settore economico (1974). (Fonte: HISItalia)

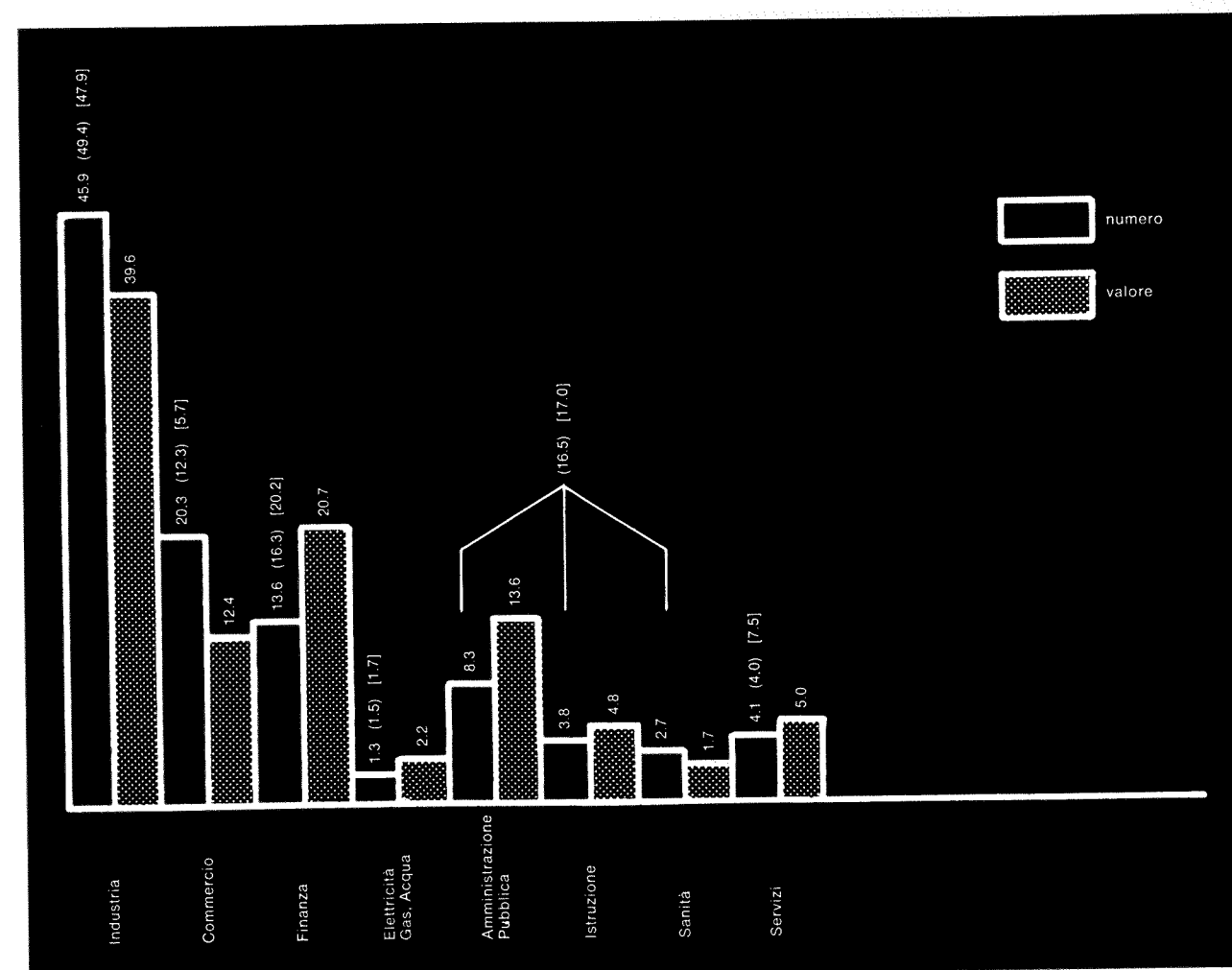


Fig. 8b) - Composizione (%) del parco calcolatori italiano per settore economico. I dati senza parentesi si riferiscono al 1974, quelli con parentesi rotonda al 1972 e quelli con parentesi quadra al 1969. (Fonte: HISItalia).

Descrizione	Valore aggiunto	Parco
Ind. alimentare	14.4	8.5
Ind. tessile	15.3	12.5
Ind. del legno	6.1	1.0
Ind. carta, editoria	5.6	6.1
Ind. metallurgica	7.5	7.5
Ind. meccanica	29.5	39.2
Ind. chimica	9.6	21.2
Ind. mat. da costruz.	7.0	2.8
Ind. varie	5.0	1.2
Tot. ind. manifatt.	100.0	34.0
Ind. Costruzioni	9.8	0.7
Ind. estrattiva, Gas, Elettr. Acqua	6.0	2.7
Commercio	16.6	12.4
Finanza	6.5	20.7
Pubblica Amministr.	12.4	20.1
Servizi vari	14.3	5.0
Totale generale	100.0	100.0

Fig. 9 - Distribuzione % valore aggiunto e valore del parco potenziale (1974). (Fonte: HISItalia).

nizzazione, che comprende, oltre all'hardware per i calcolatori general purpose, quello per i mini, i terminali ed il data entry, nonché il software, i servizi, il supporto e le spese di esercizio.

Riferendoci, sempre per il 1974, agli 8 paesi più sopra considerati, vediamo come si distribuisce la spesa complessiva per materiale EDP, che, per l'Italia, ha raggiunto nel 1974 il valore di circa 1100 miliardi di lire (vedi fig. 6). Osserviamo, innanzitutto, una estrema disparità di distribuzione della spesa fra le varie voci nei diversi paesi, riflesso ciò delle profonde disparità nell'utilizzo del calcolatore e nella sua valorizzazione. Purtroppo risaltano, per l'Italia ed il Giappone (i paesi a più recente meccanizzazione), due caratteristiche: da un lato una spesa per hardware dei calcolatori maggiore degli altri paesi e, dall'altro, una spesa per software sviluppato in house (sostanzialmente questa voce comprende gli stipendi della staff EDP). Ciò è evidente segno della relativa immaturità dei parchi calcolatori nei due paesi: infatti più bassa risulta sia la spesa per terminali (indice quest'ultimo di applicazioni maggiormente sofisticate), sia di quella per le applicazioni software acquistate dagli specialisti (servizi software e packages, ecc.).

Situata in tal modo la spesa italiana per EDP nel contesto internazionale, torniamo ad illustrare la struttura e la distribuzione del parco italiano, facendo riferimento, come più sopra abbiamo notato, al parco potenziale.

2. Struttura del parco potenziale per classe di canone mensile

Il parco italiano (fig. 7) è caratterizzato per un netto predominio dei piccolissimi e piccoli sistemi. In numero essi rappresentano l'85% circa del parco, mentre in valore la loro quota è del 36,5%. Negli ultimi 3 anni, come si vede dalla tabella, la quota (in numero) dei piccolissimi sistemi è cresciuta, passando dal 77% all'attuale valore. Questa situazione ha una razionale spiegazione: da un lato abbiamo, infatti, una struttura industriale estremamente frammentata nella quale predominano decisamente le piccole e medie imprese (da 100 a 1000 addetti, per il discorso che qui interessa) utenti potenziali di piccoli sistemi, dall'altra parte si deve considerare che la meccanizzazione, storicamente, è proceduta dalle grandi imprese scendendo verso le piccole. Fino alla metà degli anni '60, dunque, predominavano, anche in numero, i calcolatori medi e grandi, ma man mano che apparivano sulla scena nuovi utenti, di minori dimensioni, la quota dei grandi e medi sistemi andava diminuendo; sulla base dell'esperienza storica di paesi a meccanizzazione più avanzata, è possibile inferire che l'aumento della quota del numero di piccoli sistemi sul totale continuerà fino a quando i principali settori economici che utilizzano piccoli sistemi (industria e commercio) non saranno prossimi alla saturazione. Allora torneranno a prendere importanza i medi e grandi sistemi, vuoi per le progressive sostituzioni di piccoli elaboratori con altri di maggiore dimensione, vuoi per la crescita interna dello stesso parco medi-grandi sistemi. Una tale situazione è prevedibile, per l'Italia, che venga raggiunta nella prima metà degli anni '80.

3. Struttura del parco potenziale per settori economici

La struttura del parco è illustrata nei grafici di fig. 8. Dal confronto fra la distribuzione del 1974 e quella del 1969, quest'ultima presentata solo in numero di sistemi nella fig. 8 b, si deduce che l'industria è leggermente diminuita, passando dal 46.5% del totale all'attuale 45% circa; sono calate le costruzioni, mentre l'industria dell'elettricità, gas e acqua è passata dall'1.7% al 1.3%: complessivamente l'industria ha perduto un paio di punti. Anche il settore della finanza ha subito una riduzione relativa, passando dal 20% al 13.6%, mentre il commercio ha visto praticamente quadruplicata la sua presenza, passando dal 5.7% del totale al 20% del 1974: quest'ultima constatazione ci fa meglio capire lo sviluppo veramente impetuoso del mercato dei piccoli sistemi che è avvenuto in un settore come quello del commercio che presenta una struttura estremamente frammentata. Ancora diminuzioni presentano poi settori come la Pubblica Amministrazione e i servizi vari. Le cifre per il 1972 ci forniscono la fotografia della situazione intermedia.

Più interessante può sembrare la situazione presentata in fig. 9 dove è data la distribuzione percentuale

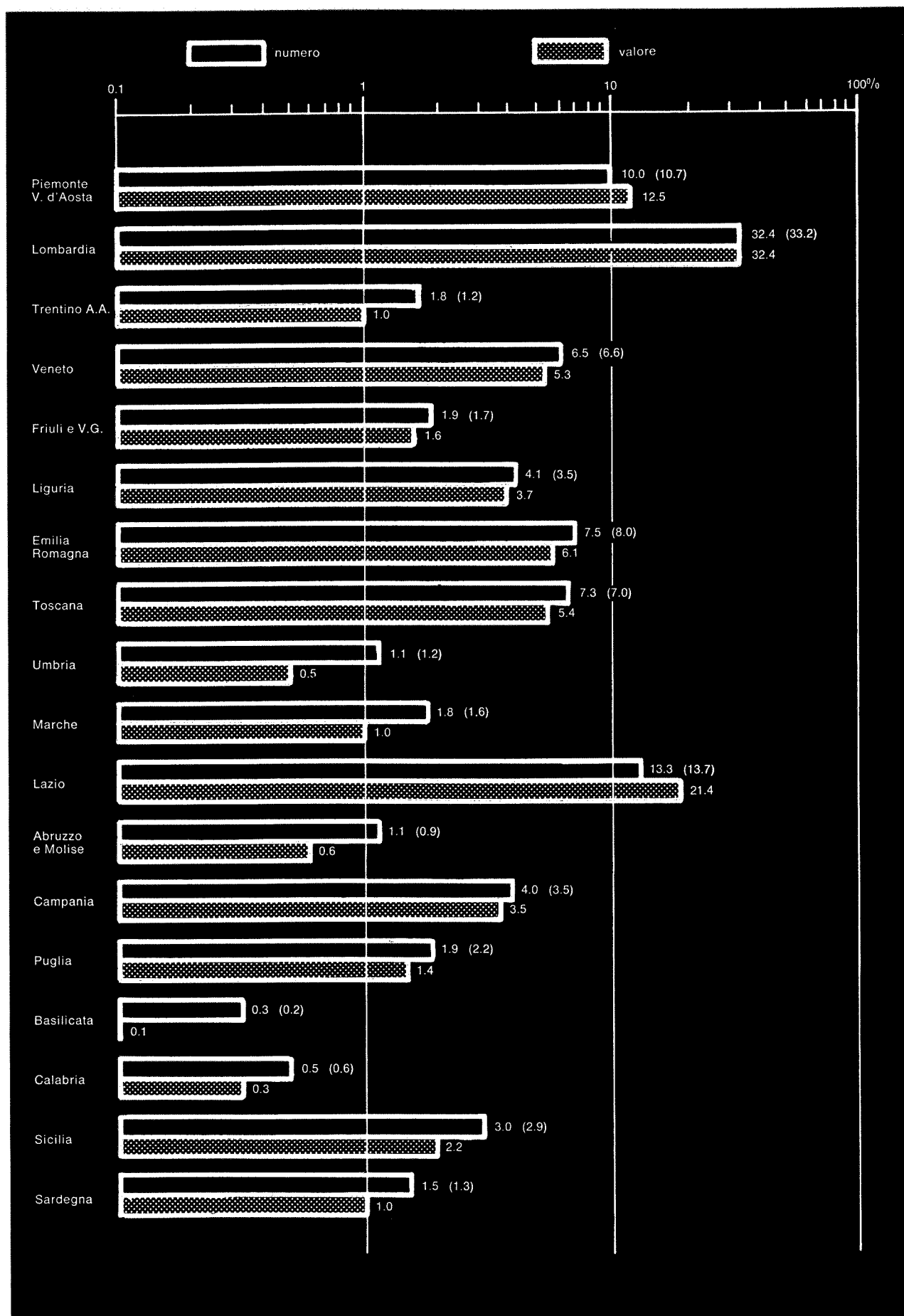


Fig. 10 - Struttura percentuale del parco potenziale italiano per regioni. I dati si riferiscono al 1974; tra parentesi, i valori del 1970. (Fonte: HISItalia).

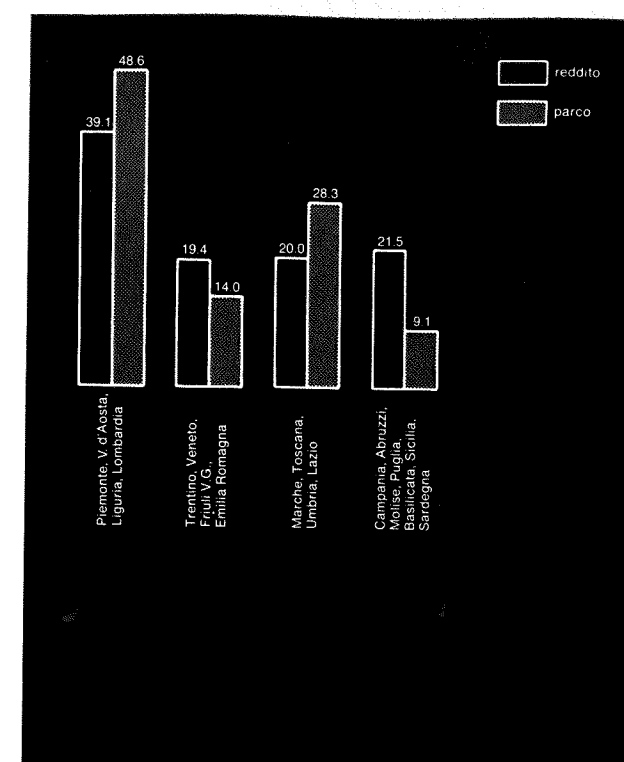


Fig. 11 - Distribuzione percentuale del reddito lordo extragricolo e del valore del parco potenziale per ripartizione geografica. (Fonte: per il reddito ISTAT, per il parco HISItalia).

del valore delle installazioni EDP nei singoli settori nel 1974 confrontata con quella del valore aggiunto sempre nello stesso anno. Si notano come alcuni settori siano particolarmente aperti alla meccanizzazione: in particolare il settore della finanza che col 6.5% del prodotto lordo assorbe circa il 21% del parco in valore e, all'interno del settore manifatturiero il settore chimico che, col 10% circa del valore aggiunto manifatturiero assorbe il 21% del valore del parco potenziale dell'industria. Squilibri di segno opposto mostrano i settori delle costruzioni e quello dei servizi: quest'ultimo però andrebbe disaggregato ulteriormente onde far risaltare il peso, al suo interno, del settore dei trasporti e comunicazioni che assorbe la quasi totalità del parco del settore « servizi ». Possiamo, dunque, schematizzare l'intera economia come caratterizzata, dal punto di vista della domanda di EDP, in due settori: il primo, di più antica meccanizzazione e di struttura economica maggiormente concentrata (meccanica, chimica, finanza) che assorbe ormai prevalentemente calcolatori di medie e grandi dimensioni, mentre altri settori, quali il commercio ed i servizi, sono caratterizzati da un ampio mercato aperto che richiede, per il momento, ancora un gran numero di piccoli sistemi.

4. Distribuzione territoriale del parco potenziale

La distribuzione territoriale del parco è fornita dalla fig. 10.

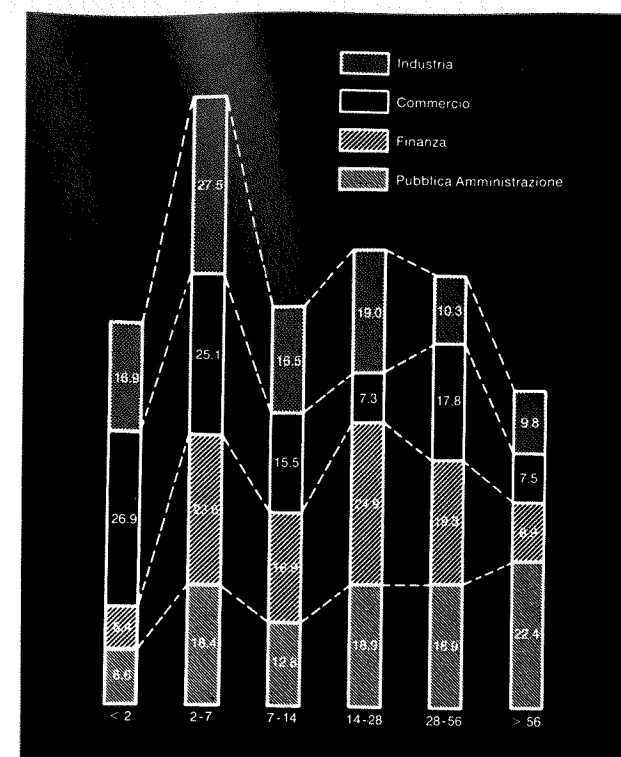


Fig. 12 - Struttura del parco potenziale in valore (1974) per i principali settori economici e per classe di ampiezza (in milioni di lire) di locazione mensile. I dati sono in % rispetto al settore economico. (Fonte: HISItalia).

Anche qui cominciamo con un confronto fra la distribuzione attuale (1974) in numero e valore e quella del 1970 in solo numero. Le variazioni nella distribuzione territoriale del numero degli elaboratori installati ed ordinati non sono state notevoli, riducendosi, sostanzialmente, a lievi aggiustamenti delle quote delle diverse regioni. Dal confronto fra distribuzione del numero dei calcolatori e quella del valore, vediamo invece, significative distorsioni. Il Lazio, ad esempio, col 13% del parco installato assorbe il 21% della spesa: l'organizzazione centrale dello Stato (Ministeri, ecc.) assorbe prevalentemente, infatti, calcolatori di grandi dimensioni; anche il Piemonte ha un parco spostato verso i grandi sistemi essendo ivi ubicati i centri di elaborazione di alcune fra le maggiori imprese nazionali. In tutte le altre regioni, il peso del numero dei sistemi è maggiore del peso dei valori, con l'eccezione della Lombardia dove c'è equilibrio.

Anche per la struttura territoriale abbiamo confrontato la struttura del reddito interno extra-agricolo con quella del valore del parco potenziale (fig. 11). Come era da aspettarsi risulta una netta prevalenza dell'Italia nord-occidentale per quanto riguarda la distribuzione del parco — prevalenza dovuta all'ubicazione delle maggiori industrie e dei più grandi istituti finanziari — e dell'Italia centrale, attribuibile questa per intero al Lazio grazie alla presenza a Roma dei principali Enti Pubblici statali e parastatali. La correlazione che esiste fra struttura economica dei settori e

Regione		Classe di ampiezza (milioni di lire)						Totale
		< 2	2-7	7-14	14-28	28-56	> 56	
Piemonte, V. d'A.	N.	53.6	25.4	9.1	8.6	3.1	0.2	100.0
	V.	17.1	17.7	14.6	27.0	20.9	2.7	100.0
Lombardia	N.	56.7	26.4	9.0	6.1	1.6	0.2	100.0
	V.	16.3	24.8	18.5	24.0	13.9	2.5	100.0
Trentino A.A.	N.	61.6	33.7	2.4	2.3	—	—	100.0
	V.	31.9	44.4	7.9	15.8	—	—	100.0
Veneto	N.	56.7	31.0	5.7	5.6	1.0	—	100.0
	V.	21.0	27.3	15.2	27.0	9.5	—	100.0
Friuli V.G.	N.	52.7	31.9	7.7	7.7	—	—	100.0
	V.	17.9	28.2	17.9	36.0	—	—	100.0
Liguria	N.	44.9	38.9	8.1	7.6	0.5	—	100.0
	V.	13.6	31.5	15.8	34.8	4.3	—	100.0
Emilia Romagna	N.	52.9	36.2	5.2	4.4	0.8	0.5	100.0
	V.	18.4	31.7	12.2	20.2	7.8	9.7	100.0
Toscana	N.	64.9	25.5	4.2	4.2	0.8	0.3	100.0
	V.	25.1	24.2	12.9	22.6	8.9	6.3	100.0
Umbria	N.	75.9	22.2	1.9	—	—	—	100.0
	V.	54.5	36.6	8.9	—	—	—	100.0
Marche	N.	69.7	27.0	1.1	2.2	—	—	100.0
	V.	39.2	38.3	5.1	17.4	—	—	100.0
Lazio	N.	41.9	29.3	9.8	13.0	4.5	1.5	100.0
	V.	7.5	14.2	11.4	29.0	24.6	13.3	100.0
Abruzzi e Molise	N.	56.9	39.2	3.9	—	—	—	100.0
	V.	29.3	56.1	14.6	—	—	—	100.0
Campania	N.	54.4	29.7	10.3	5.1	0.5	—	100.0
	V.	21.5	25.8	23.0	23.4	6.3	—	100.0
Puglia	N.	52.7	37.6	5.4	4.3	—	—	100.0
	V.	21.6	41.9	15.8	20.7	—	—	100.0
Basilicata	N.	64.3	35.7	—	—	—	—	100.0
	V.	38.2	61.8	—	—	—	—	100.0
Calabria	N.	58.3	37.5	—	4.2	—	—	100.0
	V.	29.5	37.0	—	33.5	—	—	100.0
Sicilia	N.	51.0	40.8	5.4	0.7	2.1	—	100.0
	V.	19.9	40.8	13.9	3.8	21.6	—	100.0
Sardegna	N.	67.1	21.9	9.6	1.4	—	—	100.0
	V.	32.2	31.4	29.7	6.7	—	—	100.0

Fig. 13 - Distribuzione del parco potenziale per regione e classe d'ampiezza di canone di locazione (numero e valore). (Fonte: HISItalia).

Settore	Piemonte, Liguria, Valle d'Aosta	Lombardia	Veneto, Friuli, Trentino	Emilia Romagna	Toscana, Umbria	Lazio	Resto dell'Italia	Italia
Alimentare	1.0	1.9	0.9	0.9	0.5	0.4	0.9	6.5
Tessile	1.6	3.3	1.0	0.4	1.7	0.2	0.5	8.7
Meccanica	3.3	7.2	1.9	1.6	0.7	0.7	1.2	16.7
Chimica	0.9	4.7	0.8	0.3	0.6	0.7	0.6	8.6
Resto Industria	1.1	4.3	1.4	1.0	1.4	1.3	1.4	11.9
Industria (totale)	8.0	21.4	6.0	4.2	4.9	3.3	4.6	52.4
Enti Locali	0.6	0.7	0.5	0.2	0.6	0.4	1.5	4.5
Ospedali	0.3	0.7	0.5	0.3	0.2	0.2	1.1	3.3
Banche	0.5	1.0	0.7	0.6	0.5	0.7	2.4	6.4
Commercio	3.5	8.3	3.5	2.1	2.4	3.1	3.0	25.9
Resto economia	1.3	1.3	1.0	0.5	0.5	1.7	1.2	7.5
Totale	14.2	33.4	12.2	7.9	9.1	9.4	13.8	100.0

Fig. 14 - Parco piccoli sistemi. Distribuzione percentuale per regione e settore del parco potenziale. (Fonte: HISItalia).

delle regioni e la distribuzione settoriale e territoriale del parco, non è dunque immediata: in sé il reddito prodotto non è sufficiente a spiegare le diverse propensioni alla meccanizzazione che mostrano i vari settori e le varie regioni. Ciò che è essenziale, in questo settore, è come viene prodotto il reddito e da chi: a parità di reddito industriale, infatti, mostrano maggiore assorbimento di spesa EDP quei settori maggiormente concentrati, dove cioè prevalgono le grandi imprese, o, almeno, le imprese al di sopra di una certa « soglia » che varia da settore a settore.

Ne consegue che soltanto l'analisi settoriale più puntuale può portare alla formulazione di ipotesi (per ora) volte alla spiegazione della legge della domanda di materiale EDP: ipotesi che implicano l'individuazione di un numero piuttosto elevato di parametri (quali appunto, l'ampiezza aziendale, il fatturato realizzato per le varie classi di ampiezza aziendale, il valore aggiunto, ecc.) tutti di difficile reperimento data la carenza di statistiche aziendali in Italia.

Alcuni studi, presentati sulle pagine di questa rivista [5], costituiscono un primo tentativo di esplorare sistematicamente, a livello di utilizzatore, i più importanti settori economici in Italia, al fine di poter dare, quando tutti i segmenti saranno studiati, una sintesi dell'andamento della domanda EDP insieme sufficientemente disaggregata ed empiricamente verificata.

Ma al fine di fornire qualche altro, provvisorio, elemento di giudizio sulla struttura del parco italiano di elaboratori, abbiamo provveduto ad elaborare una prima disaggregazione della struttura del parco potenziale per regioni e per alcuni dei principali settori

economici, sulla base della classe d'ampiezza degli elaboratori (fig. 12 e 13). Anche a questo primo e rozzo livello di approssimazione, alcune delle « distorsioni » più sopra evidenziate ricevono una prima, parziale, spiegazione. Ad esempio, possiamo mettere in connessione la struttura del parco potenziale del commercio, fortemente centrata sui piccolissimi elaboratori (il 27% del valore installato) con il più rapido incremento, rispetto agli altri settori, di quello del commercio: settore, come è noto, molto frammentato ed in rapida crescita nell'ultimo decennio. Dal punto di vista territoriale, poi, appare chiaro il peso che hanno nel parco del Lazio i grandi e grandissimi elaboratori, puntualmente verificato dalla struttura, in valore, del parco della Pubblica Amministrazione. Ovviamente il confronto andrebbe fatto confrontando le strutture economiche (sia dal punto di vista dei settori d'origine della produzione che da quello della classe d'ampiezza delle aziende) delle varie regioni con la distribuzione, sempre per settore di assorbimento, del parco calcolatori: l'Ufficio Ricerche di Mercato della HISItalia sta appunto completando uno studio volto a quantificare questa struttura per tutti i principali settori economici e tutte le regioni. Una anticipazione delle elaborazioni in corso può essere data nella fig. 14, ove, per gli elaboratori fino a 7 milioni di canone mensile (piccolissimi e piccoli elaboratori) viene evidenziata la distribuzione del numero di sistemi per alcuni aggregati di regioni e per i principali settori economici. Dal confronto di questa struttura con quella di appositi indicatori economici (quali l'occupazione nelle aziende o enti meccanizzabili, ecc.) dovrebbe essere possibile for-

nire, a breve scadenza, una prima sintetica « mappa » delle potenzialità di sviluppo del parco calcolatori « general purpose » per settori economici e per aree territoriali. Ciò verrebbe dunque a costituire un supporto conoscitivo non indifferente per le funzioni di pianificazione non solo all'interno delle case produttrici di materiale EDP, ma, forse prevalentemente, per quegli organi della Pubblica Amministrazione che operano nel senso di indirizzare la domanda verso obiettivi di razionalizzazione gestionale della struttura produttiva italiana.

5. Bibliografia

- [1] M. SPERANZA, *I problemi di conoscenza del mercato degli elaboratori elettronici*, in Quaderni d'Informatica, anno 1, n. 1.
- [2] *Indagine sugli elaboratori elettronici in Italia al 31 marzo 1968*, Supplemento al Bollettino di Statistica, agosto 1968.
- [3] *Produzione ed utilizzo degli elaboratori elettronici*, Congiuntura Economica Lombarda, ottobre 1969.
- [4] Comitato di Studio A.I.C.A., *La preparazione del personale per l'elaborazione elettronica dei dati in Italia*, rapporto 1971.
- [5] Vedi, ad esempio, gli articoli della dott.ssa SALA e del dott. MARANDELLA sui numeri precedenti di questa rivista.

Resoconti e recensioni

Hanno collaborato a questa rubrica:

P.L. BOSCHETTI, F. FILIPPAZZI,
L. GIANNELLI, A. MARIETTI, G. RAPELLI.

Orientamenti nella architettura dei grandi sistemi di elaborazione

U. Gagliardi, *relazione invitata al congresso « Vent'anni di informatica », Pisa 16/19 Giugno 1975*).

L'autore, un noto esperto di calcolatori (è professore ad Harvard ed ha avuto larga parte nella definizione della architettura della linea di elaboratori Honeywell Serie 60), esamina in questa memoria le nuove tendenze della architettura degli elaboratori prevedibili per gli anni '80, con particolare riguardo ai grandi sistemi.

Lo studio, partendo dalla architettura attuale dell'unità centrale e delle memorie, passa ad esaminare come tale architettura verrà influenzata nei prossimi anni dalle nuove tendenze delle tecnologie; successivamente l'indagine si estende all'evoluzione nel campo della trasmissione dei dati.

• Architettura attuale

Viene inizialmente considerata la tradizionale architettura dei sistemi attuali: un punto costante è la separazione di gestione tra la memoria principale e gli archivi del cliente residenti su memorie secondarie (dischi e nastri).

Infatti una parte del sistema operativo (Memory Management) è dedicata alla gestione degli spazi di memoria principale che vengono suddivisi tra parti residenti e transienti del sistema operativo e i vari programmi del cliente (a loro volta strutturati in parti residenti e transienti). Il Memory Management provvede al trasferimento da disco/nastro alla memoria principale delle parti transienti.

Un'altra parte del sistema operativo (Data Management) è dedicata ancora allo scambio di dati tra disco/nastro e la memoria principale, ma limitatamente alle informazioni degli archivi del cliente che, di volta in volta, sotto forma di blocchi fisici e di record logici, vengono messi a disposizione dei programmi applicativi del cliente.

Fino ad oggi la disponibilità di memorie principali di sempre maggiore capacità è stata utilizzata passando dalla mono-programmazione alla multi-programmazione e quindi con la contemporanea esecuzione di più programmi del cliente.

• Tendenze della tecnologia

L'autore esamina le tendenze delle tecnologie nel campo delle memorie principali e dei dischi:

- per le memorie principali il costo/bit è destinato a ridursi di 10 volte in 5 anni, da 0.1 a 0.01 cent/bit;

- la stessa riduzione di 10 volte in 5 anni è prevista per i dischi;

- per i nastri è prevista l'affermazione delle « librerie automatiche », in cui tutti gli archivi del cliente risultano accessibili al sistema senza intervento dell'operatore per il montaggio/smontaggio delle bobine. Con tali costi, e assumendo che il prezzo di affitto di un grande sistema rimanga circa costante (intorno ai 100.000 \$/mese) si deduce che, all'inizio degli anni '80, un tale sistema potrà disporre di:

- 128 MBytes di memoria principale

- 256 mila MBytes di capacità su disco

- 1 milione di MBytes di capacità su nastro.

I grandi sistemi avranno quindi a disposizione memorie principali e secondarie di capacità oltre dieci volte superiori alle attuali.

Questo fatto avrà come conseguenza:

- utilizzazione di dischi e nastri come archivi permanentemente in linea, per ridurre al minimo i costi e le probabilità di errore dovuti all'intervento dell'operatore per il montaggio/smontaggio;

- necessità di sfruttare diversamente la grande capacità della memoria principale, mediante una nuova architettura.

• Nuova architettura

La disponibilità di grandi capacità di memoria prin-

cipale permetterà di aumentare notevolmente il numero di programmi gestiti in multiprogrammazione. Con l'attuale architettura il numero di operazioni di lettura/scrittura (cioè accesso ad un blocco fisico su disco) per secondo, richiesto dall'insieme dei programmi, diventerebbe molto elevato e darebbe luogo ad un carico inaccettabile di lavoro per il sistema, chiamato ad eseguire l'operazione di « dispatching » (cioè trasferimento di controllo) tra un programma e l'altro, ad ogni operazione di lettura/scrittura.

Il metodo per risolvere tale problema consiste nel modificare l'interfaccia tra il calcolatore che esegue il programma del cliente (« computational processor ») e il sottosistema dischi.

L'interfaccia a basso livello (lettura/scrittura di un blocco fisico) verrà sostituita con una interfaccia a livello molto più elevato in cui la richiesta riguarderà non più i pochi records raggruppati in un blocco fisico, ma un numero molto più elevato di records tutti in relazione tra loro, secondo le complesse regole delle organizzazioni di dati di tipo data-base (record set).

Con la nuova interfaccia ad alto livello, il numero di richieste di operazioni di I/O (e di « dispatching » tra i programmi) a cui dovrà fare fronte il « computational processor » risulterà di molto ridotto.

Nasce però l'esigenza di un nuovo elemento di architettura, il « data base processor » cioè una seconda unità di calcolo che deve essere in grado di soddisfare le richieste ad alto livello (record set) del « computational processor », interpretarle, accedere ai dischi o ai nastri e trasferire l'informazione richiesta dalle memorie secondarie alla memoria principale. L'architettura di un grande sistema risulterà dunque costituita da due unità di calcolo:

— un « data base processor » con il compito di gestire la gerarchia delle memorie (nastri, dischi, memoria principale) trasferendo le informazioni dall'uno all'altro dei livelli di memoria, man mano che devono essere utilizzate dai programmi del cliente;

— un « computational processor » destinato ad eseguire i programmi del cliente, elaborando direttamente le informazioni che vengono messe a disposizione nella memoria principale da parte del « data base processor ».

Poiché le due unità di calcolo saranno completamente distinte si otterranno anche i seguenti vantaggi:

— ciascuna unità di calcolo potrà essere ottimizzata per il compito specifico che deve svolgere, con un set di istruzioni specializzato, rispettivamente, per il trasferimento delle informazioni tra un livello e l'altro di memoria oppure per l'elaborazione delle informazioni in memoria principale;

— i programmi del cliente non saranno più dipendenti dalla struttura degli archivi su nastro e su disco. Modifiche alla struttura degli archivi (per aggiunta di nuovi dati, modifiche organizzative, ecc.) non richiederanno modifiche ai programmi che elaborano tali archivi.

• Trasmissione dati

Un discorso a parte è dedicato all'evoluzione prevista per il sottosistema di trasmissione dati.

I due fattori principali che indirizzano l'evoluzione prevista sono i seguenti:

— l'introduzione, già in atto, di nuove procedure di colloquio per lo scambio di messaggi tra il calcolatore centrale e i terminali.

Le passate procedure (character-oriented) erano basate sulla suddivisione dei messaggi in caratteri di lunghezza fisica. Ciascun carattere veniva interpretato dal terminale come un dato (ad es. un numero da stampare) oppure come un carattere di controllo (ad es. inizio/fine della trasmissione, ritorno carrello, nuova linea di stampa ecc.). Si aveva così una mescolanza di dati e caratteri di controllo, un numero necessariamente ridotto dei tipi di controllo possibili e la dipendenza dal particolare codice usato per la trasmissione (5, 6, 7, 8 bit/carattere).

Le nuove procedure (bit-oriented) prevedono invece dei messaggi in cui la struttura separa chiaramente la parte di controllo (in testa al messaggio) dalla parte dati: questi ultimi non vengono considerati divisi in caratteri, ma raggruppati in un'unica sequenza di bit che non viene interpretata dalla procedura di colloquio, ma trasmessa in modo trasparente;

— disponibilità, a costi contenuti, di terminali « intelligenti », cioè programmabili e quindi in grado di eseguire delle elaborazioni di dati locali, in funzione delle particolari esigenze dell'applicazione del cliente. Gli effetti dei due fattori considerati saranno i seguenti:

— Decentralizzazione di funzioni.

Questa tendenza porta da una parte a sfruttare la disponibilità dei terminali programmabili; dall'altra parte a scaricare il calcolatore centrale di alcune funzioni specializzate (riconoscimento caratteri, trascodifica, accodamento dei messaggi, ecc.).

Nasce così un'altra unità di calcolo (Front End Processor) che si unisce ai già citati « Data base Processor » e « Computational Processor » per svolgere le funzioni un tempo svolte da un unico processor centrale non specializzato.

— Centralizzazione dello sviluppo software e della diagnostica.

Le nuove procedure « bit-oriented » sono particolarmente adatte a permettere il caricamento (via linea di trasmissione) del software e della diagnostica dei terminali intelligenti da parte del calcolatore centrale.

— « Device independence ».

Con questo termine si intende la capacità (legata alle nuove procedure « bit-oriented ») di gestire, su un'unica rete di trasmissione, terminali di tipo diverso e di poter aggiungere nuovi terminali senza modificare la rete.

Anche i programmi applicativi del cliente non sono sensibili a modifiche nel tipo dei terminali, poiché è il « Front End Processor » a gestire tali differenze; anche in tal senso i programmi applicativi risultano indipendenti dal tipo di terminale.

In conclusione l'autore prevede, nel prossimo decennio, un notevole ripensamento dell'architettura delle unità centrali dei grandi sistemi, con l'introduzione di due unità di calcolo specializzato (Computational e Data-base Processor) per adeguarsi alla nuova struttura del sistema operativo. Egli prevede inoltre nuove tendenze nella gestione delle memorie secondarie e delle linee di trasmissione.

A.M.

Corso sulla « Programmazione strutturata »

presso il CISE, Segrate (Milano), organizzato in collaborazione col capitolo italiano della ACM (Association for Computing Machinery); 2-5 dicembre 1974.

La programmazione strutturata costituisce un ampio orientamento metodologico per razionalizzare la fabbricazione del software. Su questo tema si è svolto un corso, nei giorni 2-5 dicembre 1974, presso il CISE, Centro Italiano Studi Esperienze, di Segrate. Il corso — tenuto da A. Cicu della Honeywell Information Systems Italia e da G.A. Lanzarone, M. Maiocchi e R. Polillo del Gruppo di Elettronica e Cibernetica dell'Università di Milano, con la direzione del prof. G. Degli Antoni — aveva carattere avanzato, essendo diretto a persone dotate già di esperienza in programmazione od in problemi di produzione di sistemi software.

Nel corso è stata presentata una panoramica dettagliata di tutti quegli aspetti della programmazione strutturata che sono ormai acquisiti e descritti nella letteratura specializzata; inoltre sono stati forniti risultati e valutazioni ricavati da studi ed esperienze dirette, realizzate, in particolare, presso il dipartimento di Software Engineering della HISItalia a Pregnana Milanese.

Il materiale trattato è stato raccolto in un volume di dispense, seguendo l'indice del quale si ha una traccia di quanto sviluppato nel corso. Il primo capitolo, « Introduzione alla Programmazione Strutturata », espone i concetti fondamentali di questa metodologia, ed in particolare l'uso esclusivo di strutture di controllo ad un solo ingresso e ad una sola uscita e le tecniche « top-down » di progettazione e implementazione dei programmi. Nel secondo capitolo, « Strutture di controllo », viene analizzato in dettaglio il problema delle istruzioni di controllo da impiegare, esaminando le caratteristiche dei tre costrutti fondamentali (*sequence, if-then-else, while-do*) e altre utili strutture. Il modo con cui le tecniche prima analizzate in via generale, possono essere effettivamente adottate nella pratica della programmazione sono discusse nel terzo capitolo, « Applicazione di tecniche di programmazione strutturata », che esamina separatamente i casi del COBOL, del FORTRAN e dell'assembler. Nel capitolo seguente, « Information hiding », si trattano tecniche avanzate per la costruzione di programmi facilmente modificabili e manu-

tenibili, basate sulla concentrazione di porzioni di codice che riguardano il medesimo aspetto implementativo, la cui struttura interna viene « nascosta » al resto del programma, in modo che cambiamenti rilevanti in un sistema, anche complesso, possano essere confinati all'interno di piccoli e ben definiti sottosistemi. Nel successivo capitolo, « Aspetti della comunicazione tra i componenti di un programma », viene analizzato il complesso problema delle interfacce tra i componenti di un programma o sistema software, esaminando le soluzioni tecniche che possono ridurre le difficoltà derivanti dalla scarsa modificabilità di programmi che presentano un « data-flow » complesso o, comunque, interfacce dati particolarmente articolate; in questa parte si analizza anche il problema del « recovery » e della segnalazione degli errori. Conclude il corso un capitolo su « Aspetti strutturali associati alla produzione del software », in cui si esaminano i problemi della struttura ed organizzazione di una fabbrica di software. Interessante, a questo proposito, l'impiego di opportune rappresentazioni grafiche per descrivere in modo sintetico le complesse interazioni tra le varie attività. Il corso costituisce un importante e riuscito contributo alla diffusione in Italia di queste moderne metodologie di sviluppo del software, i cui concetti saranno prossimamente oggetto di un articolo su questa rivista.

F.F.

Seminario su « Protezione e riservatezza dei dati nei sistemi informativi »

Milano, 16-17 giugno 1975, presso la Honeywell Information Systems Italia.

Su questo tema di grande attualità la Honeywell Information Systems Italia ha organizzato a Milano un seminario aperto a tutte le Aziende interessate all'argomento.

Il seminario è stato tenuto da Jerry Lobel, responsabile per la « Data Security » presso la Honeywell Information Systems a Phoenix (Arizona) e consulente in materia di numerosi enti governativi americani.

L'argomento ha suscitato grande interesse per le insospettite implicazioni, sia a carattere finanziario che a carattere legale che ne possono derivare.

La sicurezza delle informazioni affidate ad un computer è divenuta infatti negli ultimi anni un problema di rilevanti proporzioni.

Negli Stati Uniti, dove da parecchio tempo il problema è sotto esame, comitati federali appositamente istituiti hanno potuto constatare come siano in rapido aumento gli abusi, più o meno intenzionali, nell'accedere ai dati immagazzinati nei computers, con utilizzo doloso del computer per i più svariati motivi che vanno dalla concorrenza sleale, allo spionaggio industriale, alla violazione del segreto bancario, al sabotaggio.

Basti pensare che negli Stati Uniti, nel corso di una indagine, è risultato che circa l'11% delle installazioni controllate è stata vittima di uno o più accessi « criminali » al loro sistema informativo, con danni che non di rado avevano raggiunto nei singoli casi l'ordine dei milioni di dollari.

La reazione a queste scoperte è stata una proliferazione di progetti di legge atti a stroncare e a prevenire gli abusi nonché a definire giuridicamente ed a regolamentare la materia delle banche di dati.

Analoghe iniziative si ritrovano in alcuni paesi Europei come la Svezia e la Germania Federale.

Il problema della sicurezza e della riservatezza vede coinvolti sia gli utenti che le case fornitrici di hardware e software, queste ultime influenzate da motivi di prestigio, da richieste del mercato e da vere e proprie norme di tipo giuridico che condizionano in modo abbastanza rigido i concetti e le tecniche costruttive degli elaboratori.

Gli utenti, d'altro canto, dovendosi preoccupare sia di azioni dolose sia di guasti accidentali, sono alla continua ricerca di nuovi e più perfezionati metodi di sorveglianza, controllo, modalità di accesso ai computer, sistemi di allarme, possibilità di « back-up » e così via.

Anche i programmi applicativi e il progetto dei sistemi informativi subiscono modifiche per far fronte a questa nuova esigenza.

L'organizzazione stessa delle Aziende è coinvolta in questo processo di adeguamento e cerca di far fronte a queste nuove problematiche con l'istituzione di commissioni di « auditing » che eseguono improvvisi controlli, la ricerca di metodi di selezione e di indagine particolari per il reclutamento del personale da utilizzare in ambienti EDP, l'estensione dell'automazione a tutti i settori in cui è applicabile, per limitare al massimo il fattore umano.

Da questa breve panoramica si può intravedere la vastità delle implicazioni che il problema della sicurezza e della protezione delle informazioni si trascina dietro nella non facile ricerca della soluzione ideale.

La Honeywell Information Systems ha fatto di tale problematica e delle relative soluzioni un obiettivo di fondo della propria attività di ricerca e sviluppo specialmente nel settore dei grandi sistemi, e i risultati già ottenuti la pongono all'avanguardia in tale campo.

Il sistema Multics, realizzato dalla Honeywell in collaborazione con il Massachusetts Institute of Technology (MIT), e recentemente messo sul mercato come Livello 68 della nuova serie 60, è stato infatti giudicato dalla Commissione del Senato americano che si occupa della « privacy », come il sistema che oggi meglio garantisce da ogni intrusione non voluta nelle informazioni in esso contenute.

La sempre maggior esigenza da parte delle utenze di sistemi informativi di un servizio continuato e sicuro, il sempre più vasto impiego delle banche dei dati, a cui vengono affidate informazioni molto riservate, fanno sì che il tema della sicurezza e della

riservatezza rappresenti in sempre maggior misura l'obiettivo da raggiungere, sovrapponendosi alla precedente tendenza che richiedeva agli elaboratori unicamente maggior potenza di calcolo.

L.G.

Statistica ed elaborazione statistica delle informazioni

di Fausto Ricci, pagg. 383, Società Zanichelli, 1974.

L'autore, che fa parte della Honeywell Information Systems Italia ed insegna all'Università di Milano, tratta in questo volume un tema di particolare importanza nel quadro della elaborazione delle informazioni.

Il tema principale svolto è l'analisi della variabilità ad una o più dimensioni, sulla base delle manifestazioni osservate di un fenomeno, dette informazioni elementari. Nella prima parte, essenzialmente descrittiva, le informazioni riguardano una sola componente e vengono sintetizzate con indici classici quali la media, la varianza, il rapporto di concentrazione. Vengono anche esaminate le proprietà delle principali variabili casuali, tipiche strutture di riferimento secondo cui interpretare le caratteristiche delle distribuzioni ad una dimensione.

Nella seconda parte sono esaminati fenomeni con due componenti e sono presentati alcuni strumenti per valutarne il grado di correlazione e di connessione. Gli stesi concetti sono poi generalizzati nella parte quinta, dove vengono proposte due classi di modelli interpretativi di sistemi complessi.

Si tratta della teoria della regressione multipla e delle componenti principali, che tanta parte hanno nella spiegazione della variabilità di fenomeni socio-economici, quali la risposta ai tests attitudinali ed il comportamento del mercato dei consumatori.

Nella parte terza si studia la dinamica dei sistemi, con particolare riferimento alla teoria dei processi Markoviani, le cui tipiche applicazioni si hanno nell'ambito delle file d'attesa, del cambiamento di atteggiamenti, ecc.

Vengono infine trattate la teoria della stima e della verifica dell'ipotesi, analizzando le caratteristiche dei test di significatività classici e moderni.

Nelle appendici sono riportate l'elenco dei simboli e delle abbreviazioni usate, alcune funzioni e disuguaglianze notevoli, le tavole numeriche di uso più comune in Statistica e Ricerca Operativa, una bibliografia selezionata e l'indice analitico, seguendo una impostazione che agevola la consultazione del testo, la cui comprensione è facilitata da un gran numero di figure e problemi svolti.

Le formule sono presentate in un modo adatto all'applicazione su elaboratori, seguendo il principio di minimizzare a priori il costo di elaborazione delle informazioni.

Il linguaggio è rigoroso ma abbastanza comprensibile sia agli studenti universitari che ai tecnici di uffici

studi, ricerche, ed in generale analisti EDP: si tratta in sostanza di un ulteriore ponte fra università ed industria.

P.L.B.

Scheutz - La macchina alle differenze

(Un secolo di calcolo automatico)

a cura di Mario G. Losano, pagg. 164, ETAS Libri 1975

Dopo il « Babbage » dello scorso anno, Losano dedica questo nuovo volume della serie « Un secolo di calcolo automatico » a un personaggio che senza dubbio, nei confronti del matematico inglese, rappresenta una figura « minore » ma che pure merita una collocazione non marginale nella storia delle macchine da calcolo: lo svedese Georg Scheutz.

Proprio a Scheutz, e a suo figlio Edvard, si deve infatti la prima compiuta realizzazione di una « macchina alle differenze », ossia di una macchina addizionatrice stampante a più stadi, capace di sviluppare funzioni e quindi di calcolare tavole dei più vari tipi (astronomiche, nautiche, attuariali ecc.) con il metodo delle differenze.

La prima idea di una macchina di questo tipo era venuta a Babbage, come egli stesso ricorda, nel 1812 o 13 ma il progetto relativo prese corpo solo negli anni 20 e si trascinò a lungo, la prima macchina differenziale di Babbage venendo presentata solo nel 1862 alla Esposizione di Londra.

Nel frattempo però Babbage ne aveva scritto nella sua opera « Economy of Manufactures and Machinery » uscita nel 1832 e una descrizione della macchina, o meglio del suo progetto, era stata data da Lardner nel numero di luglio 1834 della « Edinburgh Review ».

Entrambi questi scritti erano noti a Scheutz e, abbia tratto lo spunto dall'uno o dall'altro, fatto sta che alla metà degli anni 30 egli, che allora aveva cinquant'anni e faceva il tipografo e l'editore di riviste tecniche a Stoccolma, cominciò a lavorare, aiutato dal figlio Edvard appena quindicenne, a un suo modello di macchina differenziale.

Come ebbe a riconoscere successivamente lo stesso Babbage, la macchina di Scheutz ha in comune con la sua ben poco all'infuori del principio ma è frutto di un lavoro di sviluppo quasi interamente originale. Nel libro sono raccontate dettagliatamente le difficoltà, soprattutto finanziarie, che gli Scheutz dovettero superare per arrivare alla costruzione della macchina. Fin dal 1838 Georg Scheutz aveva chiesto un

finanziamento al governo svedese ma solo nel 1851 gli venne concesso un contributo modestissimo (meno del 4% di quello che era stato speso fino ad allora in Inghilterra per l'analoga macchina di Babbage) e a condizione che la macchina fosse presentata funzionante entro un termine piuttosto breve. La macchina di Scheutz venne pronta nell'ottobre 1853, poteva operare con numeri di 15 cifre e svolgere funzioni tabulari con quattro ordini di differenze di cui la quarta costante. I risultati venivano stampati su strisce di cartapesta spalmate di piombo atte a ricavarne direttamente lastre per la stampa in stereotipia delle tavole (evitando così ogni trascrizione). Nel 1855 la macchina venne portata a Londra ed esposta alla Royal Society. Nello stesso anno raggiunse l'Esposizione di Parigi dove ottenne la medaglia d'oro. Altri riconoscimenti vennero agli Scheutz grazie soprattutto agli sforzi di Babbage che, va detto a suo onore, anziché ingelosirsi si era fatto il loro più acceso sostenitore.

Il libro narra le vicende successive di questa macchina (che finì in America, acquistata dall'osservatorio astronomico di Albany, e ora si trova allo Smithsonian di Washington) e della seconda analoga, realizzata da Edvard Scheutz e Bryan Donkin nel 1856 e acquistata dal governo inglese per il General Registrar Office che se ne servì per la preparazione delle tavole di vita della popolazione inglese (questa seconda macchina si trova ora al Museo della Scienza di Londra).

Il libro accenna anche ad alcune altre macchine alle differenze realizzate in Svezia e in America nella seconda metà dell'800. Il filone si esaurì peraltro abbastanza presto, le macchine alle differenze — che erano macchine destinate a un impiego specifico — venendo soppiantate da macchine da calcolo destinate a un uso più generale.

La parte del libro dedicata alle vicende della macchina di Scheutz e delle altre macchine differenziali è opera di Uta Merzbach, responsabile della Sezione Matematica dello Smithsonian mentre è di Losano il saggio biografico introduttivo. Una estrema cura filologica caratterizza entrambe le parti e, nella grande scarsità almeno in Italia di opere relative alla storia del calcolo automatico, fa rimpiangere che un tale sforzo non venga dedicato a imprese di più ampio respiro.

Completano il libro una parte antologica (qui peraltro assai meno interessante, e pour cause, che non l'analoga parte del volume su Babbage) e una serie di pregevoli illustrazioni.

G.R.