

Quaderni di informatica 21

Rassegna di tecnologia ed applicazioni degli elaboratori elettronici - Anno X - Numero 3 - 1983



Honeywell
Honeywell Information Systems Italia

Quaderni di informatica

Rassegna di tecnologia ed applicazioni degli elaboratori elettronici
Anno X - Numero 3 - 1983

Sommario

Tecnologie al setaccio

Nuovi dispositivi logici e di memoria sono stati e vengono continuamente proposti. Pochi, però, arrivano ad essere utilizzati praticamente.

F. FILIPPAZZI

Comunicazione mediante fibre ottiche

I sistemi di comunicazione con fibre ottiche costituiranno fra non molto il mezzo più economico e flessibile per molte applicazioni, in particolare per realizzare reti informatiche locali.

A. HUSAIN, S.D. COOK

Tecniche di individuazione degli errori nella trasmissione dati

Viene fatta una panoramica delle tecniche impiegate per rilevare gli eventuali errori introdotti nel processo di trasmissione dei dati.

T. HOUSLEY

L'elaboratore nello stoccaggio di gas naturali in pozzi esauriti

Viene presentata una interessante soluzione del problema di realizzare serbatoi polmone per gas naturale. Il computer è un elemento essenziale della soluzione.

E. FIAMBERTI

Macchine calcolatrici e intelligenza

Viene qui proposto, in chiave storica, uno scritto di uno dei maggiori geni della logica matematica, fondatore della teoria delle macchine calcolatrici.

A.M. TURING

Direttore Responsabile
F. Filippazzi

Comitato di Redazione
A. Cicu, C. Falcetti, P. Lupo,
M. Nobile, G. Occhini
G. Rapelli, F. Sala

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

Stampa
tipolito Maggioni s.a.s.

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
esprese rimane agli Autori.

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

Tecnologie al setaccio

FRANCO FILIPPAZZI

*Honeywell Information Systems Italia
Milano*

1. Introduzione

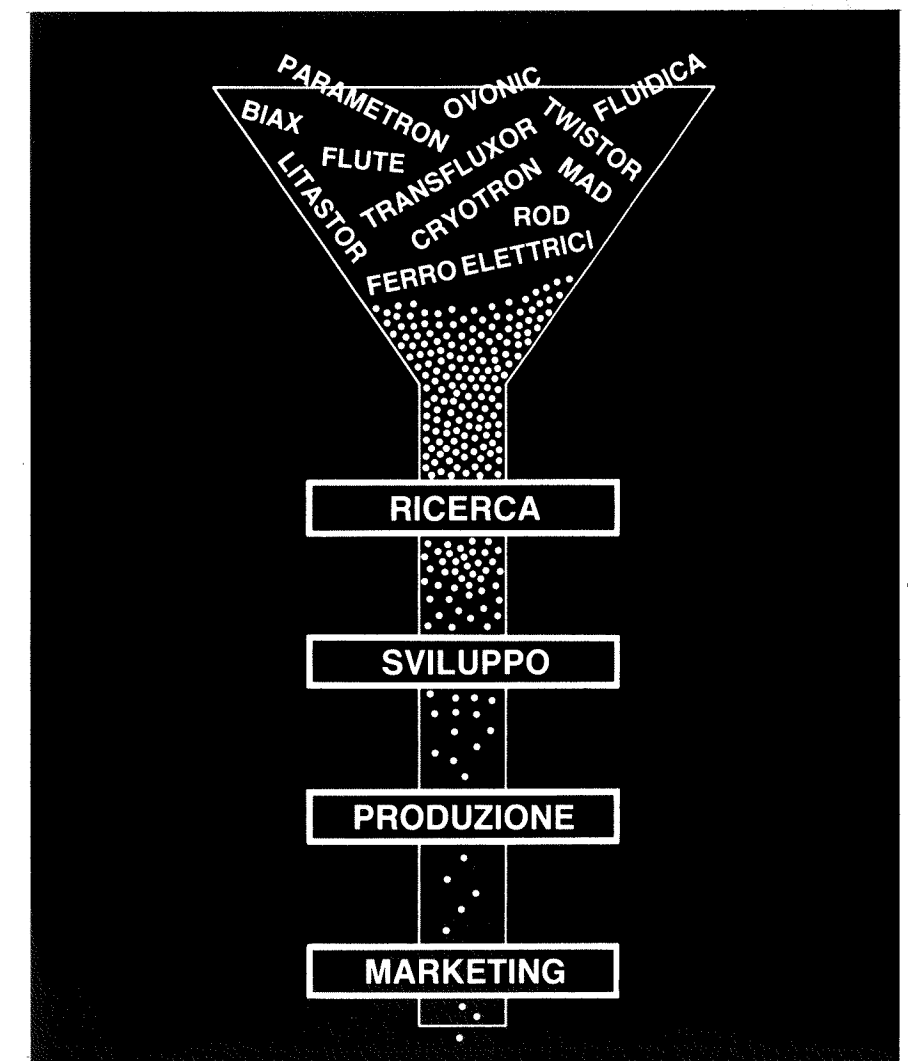
È un fenomeno fisiologico della ricerca in settori avanzati che solo una piccola parte delle innovazioni proposte raggiunga la fase di effettiva utilizzazione (fig. 1). Molte idee

muoiono quietamente, a vari stadi dello sviluppo, nell'ambiente dove sono nate, senza che il mondo esterno ne venga nemmeno a conoscenza.

Accade però che nuove soluzioni

vengano clamorosamente alla ribalta, per poi svanire più o meno rapidamente nel nulla, senza neanche un epitaffio che indichi le cause della loro scomparsa. Forse si trattava di un "baco" fatale, rivelatosi però

Fig. 1 - Il setaccio delle idee. È fisiologico che gran parte delle innovazioni proposte si fermino strada facendo in corrispondenza dei diversi filtri. Quelle che escono dal setaccio diventano prodotti industriali.



solo dopo che l'annuncio del "rivoluzionario dispositivo" era stato fatto. O forse il difetto era noto ai ricercatori, che ritenevano però di porvi rimedio prima della messa in produzione. O forse ancora il dispositivo funzionava regolarmente, ma non era producibile ai costi previsti. O forse nel frattempo era nata una tecnologia più competitiva. O forse...

Si potrebbe continuare con le ipotesi, ma è più istruttivo riferirsi a casi reali. E allora, per quanto riguarda i computer, c'è solo l'imbarazzo della scelta. Si potrebbe infatti fare una lista lunga qualche pagina semplicemente coi nomi dei dispositivi che hanno avuto un più o meno lungo momento di celebrità e sono poi scomparsi dalla scena.

2. La logica magnetica

Prima dell'avvento dei circuiti integrati a semiconduttore si ebbe un grande fervore nella ricerca di soluzioni che minimizzassero il numero degli elementi attivi dei circuiti (tubi prima, transistori poi).

A tale scopo, si cercò in particolare di realizzare le funzioni logiche richieste mediante dispositivi magnetici di forma appropriata.

Venne a tal fine ideata tutta una serie di dispositivi chiamati genericamente MAD (Multi-Apertured Devices), ossia elementi magnetici a molte aperture. Il loro funzionamento si basava sul fatto che mediante correnti inviate attraverso alcuni dei fori è possibile controllare i segnali generati nei fili passanti per altri fori.

Il capostipite di questi elementi fu il cosiddetto TRANSFLUXOR, a cui seguirono molti altri, alcuni illustrati in fig. 2, con cui si realizzavano funzioni booleane più o meno complesse.

Una soluzione particolare di logica magnetica fu il PARAMETRON, costituito sostanzialmente da un circuito risonante con due stati stabili di oscillazione, sfasati di 180°, mediante i quali veniva rappresentata l'informazione binaria. Su questo principio fu costruito in Giappone agli inizi degli anni '60 un intero elaboratore, poi abbandonato.

Si può ulteriormente citare, tra le soluzioni proposte, la logica a ROD ("barre"), i cui elementi costitutivi erano appunto delle sottili barre di vetro ricoperte di materiale magnetico, ove, tramite opportuni avvolgi-

menti, venivano create le richieste configurazioni di flusso.

Questi ed altri sistemi di logica magnetica costituirono per anni un notevole campo di ricerca e sviluppo. Parecchi dispositivi vennero prodotti e messi sul mercato; però, in pratica, il loro impiego non andò oltre la costruzione di prototipi o di qualche apparato particolare.

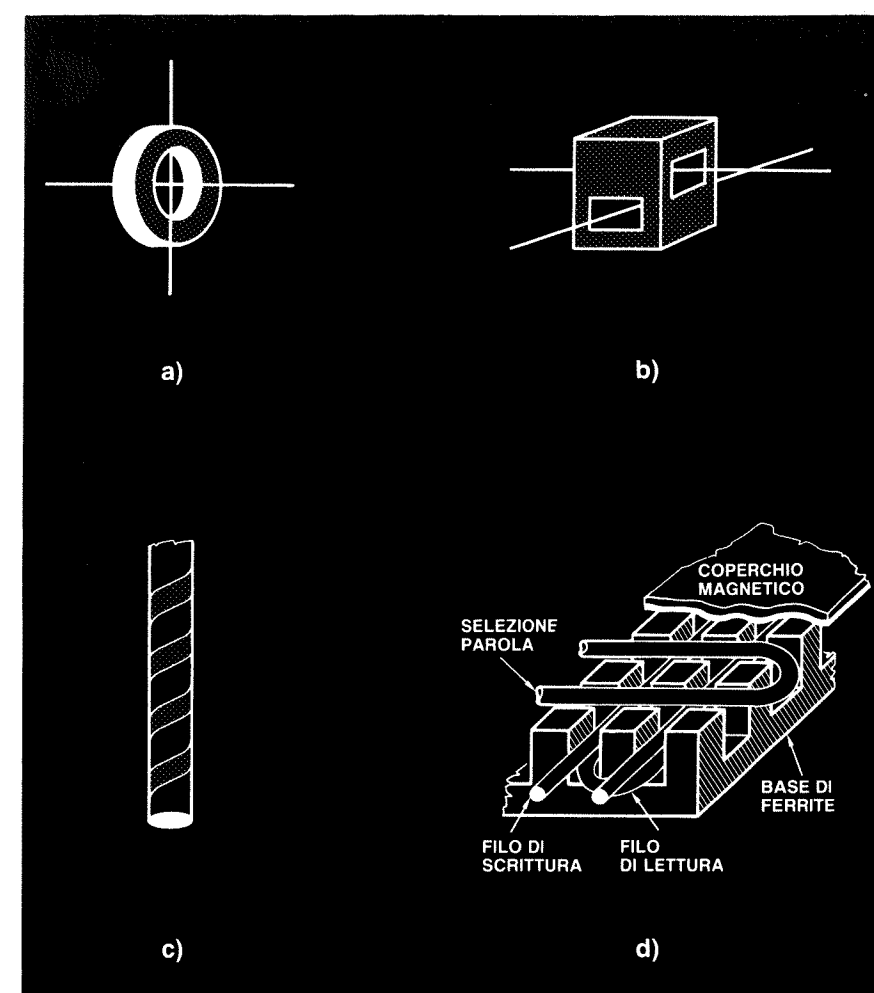
Quali le ragioni dell'insuccesso? In sostanza, le prestazioni ottenute con la logica magnetica non raggiunsero mai quelle dei circuiti logici a valvole o a transistori; successivamente, l'avvento della tecnologia dei circuiti integrati eliminò completamente ogni possibilità di competizione.

3. Quando vince il "toro"

La tecnologia di memoria è stata dominata dalla metà degli anni '50 fino agli inizi degli anni '70 dall'anellino ("toro") di ferrite. Durante tale periodo, furono proposte decine di soluzioni alternative, aventi l'obiettivo sia di migliorare le prestazioni delle memorie toroidali sia di superarne le difficoltà costruttive. Come si ricorderà, gli anellini (di diametro

Fig. 3 - Il nucleo toroidale (a), elemento costitutivo delle memorie centrali per circa tre lustri, e alcuni dei suoi competitori: (b) il BIAX, elemento con due assi, che consente una interrogazione non distruttiva, (c) il TWISTOR, qui mostrato nella sua forma originaria a "insegna di barbiere" (barber-pole) e (d) il WAFFLE IRON, struttura a "stampo per cialde". Queste ed altre soluzioni sono state ideate in particolare per facilitare l'assemblaggio della memoria.

(Per ulteriori informazioni si veda: W. Reinwick, A.J. Cole, *Digital storage systems*, Chapman & Hall, 1971)



esterno attorno al millimetro) dovevano essere montati infilandovi diversi fili di rame, e "cuciti" faticosamente tra loro in matrici contenenti migliaia di elementi.

Le alternative più studiate partivano ancora da materiali magnetici, ma usavano geometrie e meccanismi di funzionamento differenti (fig. 3).

Un dispositivo da citare è il BIAX, un blocchetto di ferrite con due fori tra loro ortogonali. Questo elemento consentiva, in particolare, la lettura dell'informazione senza implicarne la cancellazione, come invece avveniva nei nuclei toroidali.

Una soluzione a suo tempo celebre fu il TWISTOR, ideato nei laboratori Bell, costituito nella sua forma originaria (chiamata "Barber-pole", per la somiglianza con l'insegna dei barbiere negli Stati Uniti) da un nastro di permalloy avvolto a spirale su un supporto cilindrico.

Furono ideate memorie dalle forme più diverse, come il FLUTE ("flauto"), un lingotto di ferrite con dei fili

annegati; il WAFFLE IRON ("stampo per cialde") costituito da blocchi ferritici scanalati per farvi passare i fili; i FERRITE BEAN ("fagioli di ferrite"), costituiti da minuscole palline di impasto magnetico collocate agli incroci di una matrice di fili, e cotte poi in un forno; e così via.

Malgrado questo fiorire di idee, il nucleo toroidale continuò in pratica a dominare incontrastato il campo delle memorie veloci. Giocarono a suo favore proprio la sua semplicità e conseguentemente la facilità di produzione e collaudo, e in definitiva il suo miglior rapporto costo/prestazioni. Tutti gli altri dispositivi ebbero il loro momento di gloria; ma di fatto non lasciarono tracce. È il caso di dire che in questa arena stavolta ha vinto il toro...

4. Il computer che non venne dal freddo

Nel 1956, D.A. Buck, un ricercatore del MIT, il famoso Massachusetts Institute of Technology, propose un

nuovo e rivoluzionario modo di costruire gli elaboratori.

Il dispositivo di base, da lui chiamato CRYOTRON, si basava sul fenomeno della superconduttività, ossia sull'annullamento della resistenza elettrica presentato da certi materiali a temperature prossime allo zero assoluto (-273°C).

Per quanto potesse sembrare fantascientifica, la poposta aveva dei ragionevoli fondamenti. Un computer a cryotron presentava infatti vantaggi potenziali, sia funzionali che economici, decisamente superiori a quelli delle altre tecniche dell'epoca e tali da compensare il fatto che l'ambiente di funzionamento fosse del tutto fuori del comune. (Tutta la macchina doveva lavorare immersa in un bagno di elio liquido, a circa 2 gradi assoluti).

L'idea di Buck ebbe una ampia evoluzione tecnologica e costruttiva (fig. 4). Dal dispositivo originario con fili avvolti, si passò a soluzioni a strati sottili, depositati su lastrine di

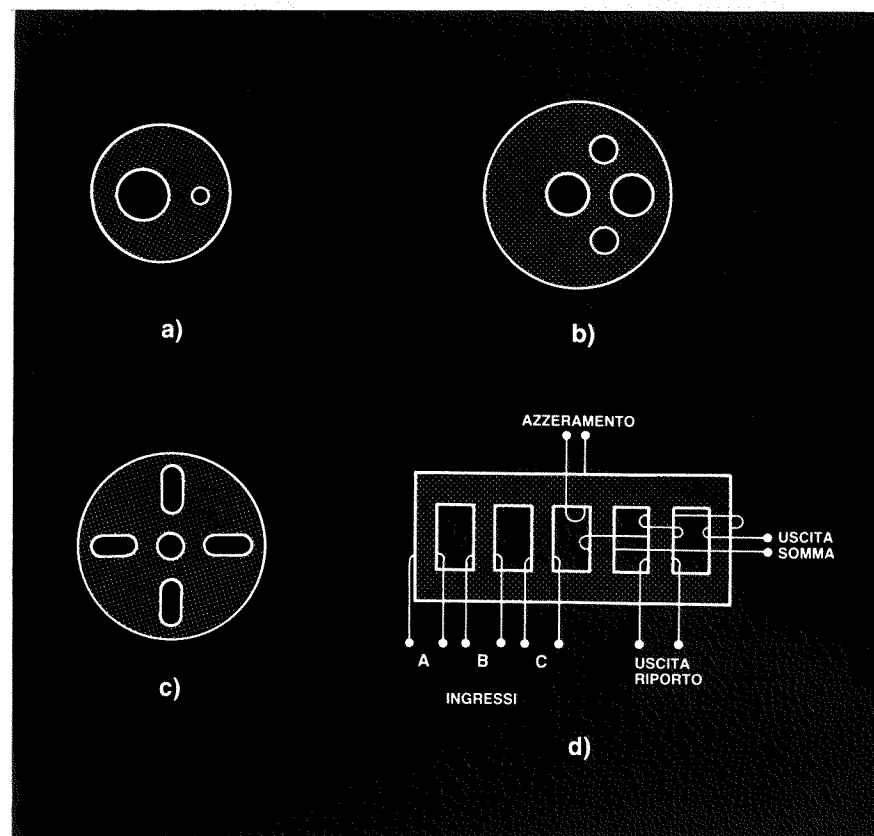


Fig. 2 - Realizzazione di funzioni logiche mediante dispositivi magnetici a più aperture (MAD). Inviando correnti in alcune delle aperture si generano segnali di uscita su altre. È così possibile realizzare funzioni logiche abbastanza complesse in un unico dispositivo. In figura, (a) rappresenta l'elemento più semplice originario (TRANSFLUXOR), (b) realizza un OR logico, (c) un AND logico, e (d) un sommatore di tre cifre binarie.

(Per ulteriori informazioni si veda: A.J. Meyerhoff, *Digital applications of magnetic devices*, J. Wiley, 1960)

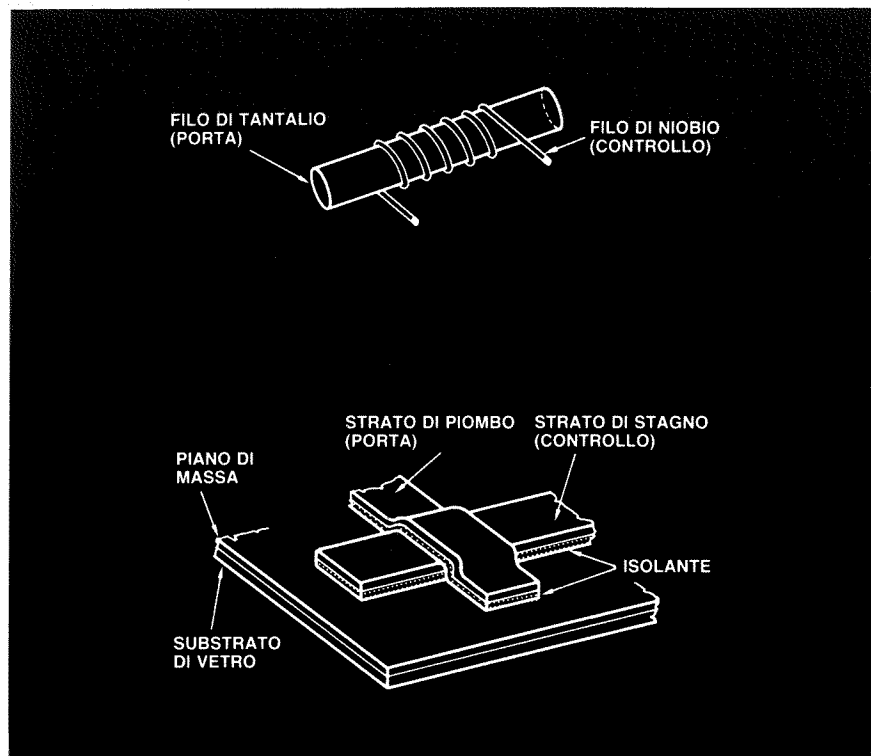


Fig. 4 - Risale al 1956 l'idea di realizzare computer operanti a temperature vicine allo zero assoluto, per sfruttare il fenomeno della superconduttività. La fig. (a) mostra il primo dispositivo proposto (CRYOTRON), nella sua forma originaria. Il filo di tantalio funziona da "porta"; il suo stato (resistivo oppure superconduttore) viene controllato dalla corrente inviata nel filo di niobio. La fig. (b) mostra una versione successiva del CRYOTRON, realizzata con strati sottili depositati su un substrato di vetro.

(L'articolo originale di D.A. Buck sul cryotron è apparso su "Proceedings of I.R.E.", vol. 44, 1956).

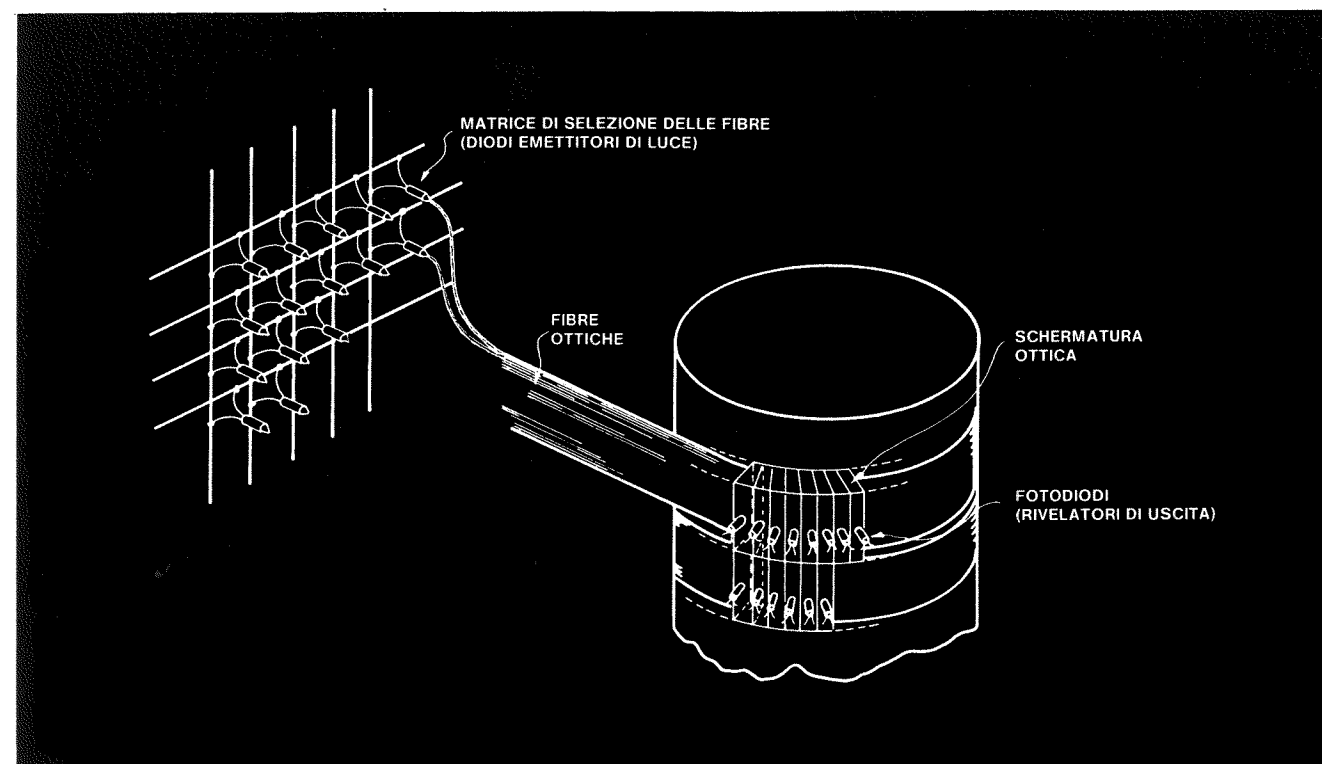


Fig. 5 - Il LITASTOR è un esempio di impiego della luce per realizzare memorie non cancellabili ("Read-Only") di grande velocità.

Le informazioni sono registrate lungo fibre ottiche sotto forma di microincisioni, che costituiscono altrettanti punti di "spillamento" della luce inviata nella fibra. Una informazione binaria è rappresentata dalla presenza o meno di una microincisione nella posizione assegnata.

(Articolo originale: F. Filippazzi, *A new approach to permanent memory*, IEEE Trans. on Electronic Computers, Vol. 16,3,1967).

vetro, che consentivano elevati livelli di integrazione e miniaturizzazione. Diverse aziende, come General Electric e IBM, effettuarono ampi investimenti di ricerca nel settore. Tutti questi sviluppi vennero in pratica abbandonati nel periodo attorno alla metà degli anni '60. Cosa era successo? Semplicemente che la tecnologia dei semiconduttori era ormai esplosa e si presentava come la soluzione vincente nel campo della microelettronica.

Non vi era perciò ormai nessuna ragione per perseguire una tecnologia che implicava condizioni ambientali estreme, quando risultati paragonabili potevano essere ottenuti con circuiti operanti nel normale campo di temperatura. Pertanto di cryotron, da allora, non si sentì più parlare.

In effetti, l'idea del computer criogenico venne ripresa verso la fine degli anni '60, basandosi però su un nuovo e diverso fenomeno delle bassissime temperature (l'effetto Josephson), che promette velocità di funzionamento elevatissime. Sono tuttavia passati quasi quindici anni da allora, e il "calcolatore che viene dal freddo" non è ancora in vista.

5. Fotoni contro elettroni

Quando si parla di elaboratore, si

sottointende che esso è "elettronico".

Questa definizione viene messa in discussione dai tentativi diretti a realizzare computer il cui funzionamento sia basato su fenomeni ottici anziché elettrici, per sfruttare le peculiari caratteristiche della luce, in particolare la sua velocità.

L'idea di usare la luce nei calcolatori è tutt'altro che nuova e possiamo anche in questo caso fare un notevole elenco di soluzioni abortite o scomparse in tenera età.

A cominciare dalla metà degli anni '60 molti sforzi sono stati fatti, in particolare, per realizzare memorie ottiche, che costituissero un salto in avanti rispetto alle soluzioni magnetiche allora imperanti (nuclei toroidali per le memorie centrali, registrazione su superfici magnetiche per le memorie di massa).

Memorie ottiche a punto di Curie, dispositivi a effetto Faraday, memorie fotocromatiche, memorie a linee di ritardo ottiche, ecc. sono nomi noti agli addetti ai lavori, che non hanno però portato al di là di prototipi di laboratorio.

Se ci è consentito un inserto personale, si può citare tra gli altri il LITASTOR ("Light Tapping Store", ossia memoria a spillamento di lu-

ce), proposta dallo scrivente nel 1967 e basata su un impiego inconsueto delle fibre ottiche, allora ai primordi (fig. 5).

In sostanza, nel campo della elaborazione ottica si sono registrate molte idee, si è generata un'ampia letteratura scientifica, ma non c'è stato nessun vero prodotto industriale.

Una ragione di fondo di questo stato di cose è che certamente gran parte delle soluzioni proposte non offrivano i livelli di affidabilità, riproducibilità e costo richiesti per farne un oggetto commerciabile.

Senza contare che nel frattempo le tecnologie già in uso rimanevano tutt'altro che ferme, costituendo un "bersaglio mobile" per quelle in competizione. Basti pensare alla tecnica della registrazione magnetica, in cui la densità di informazioni è continuata ad aumentare costantemente da vent'anni a questa parte, decuplicando il suo valore circa ogni cinque anni.

Non si può chiudere il capitolo delle applicazioni della luce negli elaboratori senza menzionare un caso clamoroso e sotto certi aspetti divertente. Nel 1971 il mondo dei computer venne messo in agitazione da Frank Marchuk col suo "computer a laser". Doveva essere questo un elabo-

ratore con ciclo operativo di 20 nanosecondi e una memoria di 10 trilioni di bit. Lo strabiliante CG-100 (questo il nome della macchina) venne mostrato a molti potenziali acquirenti (tra cui il governo degli Stati Uniti e le più grandi aziende di elaboratori), ma mai in funzionamento.

Alcuni mesi dopo l'annuncio, un gruppo di esperti del CALTEC (California Institute of Technology) mostrarono come alcuni dettagli del CG-100 rappresentassero una chiara violazione del principio di indeterminazione di Heisenberg, e che neanche col laser si sarebbe mai potuto realizzare una memoria così densa come veniva affermato.

E così, dopo essere comparso persino sul "Wall Street Journal" ed aver fatto per qualche mese i titoli della stampa specializzata, di Marchuk e del suo "laser computer" (pronunciato poi ironicamente "leisure computer") non si seppe più nulla.

Quello di Marchuk è evidentemente un caso ai limiti della truffa e non va confuso con tutti gli altri esempi cui si riferisce questo scritto. Dietro ognuno di essi c'è infatti, di norma, un lungo e serio lavoro di ricerca scientifica e tecnologica, e molto spesso il lampo di genio e l'intuizione creativa.

6. E l'elenco continua...

Potremmo aggiungere al nostro elenco ancora molti nomi, più o meno esotici, come OVONIC (dispositivo a semiconduttori amorfi), PHOTOSTORE (memoria a disco fotografico), CHARACTERON (sistema di stampa da tubo a raggi catodici), MEMORIA FERROELETTRICA (duale delle comuni memorie ferromagnetiche), LOGICA FLUIDICA, e così via.

Sono però nomi che, se dicono molto a chi ha vissuto da addetto ai lavori la tumultuosa evoluzione dell'elaboratore, sarebbero poco significativi per il lettore generico.

Comunicazione mediante fibre ottiche

A. HUSAIN
S. D. COOK

Honeywell Inc.
Technology Strategy Center
Minneapolis (USA)

1. Introduzione

Le fibre ottiche sono un mezzo di comunicazione in cui la luce è trasmessa all'interno di una fibra sottile e flessibile di vetro o di plastica. Rispetto alle prime esperienze degli anni '50, si sono fatti enormi progressi. Attualmente vengono ottenute fibre con attenuazione al di sotto di 0,4 dB/km e con larghezza di banda di alcuni GHz/km. L'uso più comune è oggi quello di un mezzo di trasmissione che connette due unità o circuiti elettronici. Il collegamento mediante fibra ottica converte in luce un segnale elettrico, trasmette la luce attraverso la fibra e riconverte la luce in un segnale elettrico. Esso converte gli elettroni in fotoni e viceversa. Fondamentalmente, questo comporta una sorgente di luce, come un diodo emittitore di luce o un laser, e un rivelatore di luce, come un fotodiodo.

Fino a poco tempo fa, la principale applicazione della tecnologia delle fibre ottiche è stata per comunicazioni punto a punto e per limitate applicazioni militari. Prestazioni superiori e costi in rapido declino dei dispositivi a fibre ottiche, insieme ai vantaggi inerenti a tale supporto, stanno aprendo nuove aree di applicazione come le reti locali (LAN), circuiti di abbonati locali e distribuzione a breve distanza di dati e di segnali video. In effetti, il mercato dei dispositivi a fibre ottiche è previsto nella fascia di 1-2 miliardi di dollari per il 1990. Questo articolo intende presentare i vantaggi e gli elementi

base di un sistema a fibre ottiche e fornire un breve panorama dello stato delle comunicazioni mediante queste fibre.

La fibra ottica è fisicamente piccola, flessibile, robusta, leggera ed economica; da un punto di vista ottico, possiede una grande larghezza di banda per trasportare informazioni, ha basse perdite e proprietà particolari che permettono di sopportare variazioni di temperatura estreme, resistere all'interferenza di radiazioni elettromagnetiche, non dare luogo a interferenza elettromagnetica e non essere causa di incendio. Queste caratteristiche distinguono le fibre ottiche dai sistemi convenzionali.

Larghezza di banda superiore: la capacità potenziale di trasportare informazioni cresce con la frequenza. Poiché le frequenze della luce sono di alcuni ordini di grandezza superiori alle frequenze radio, una fibra ottica ha il potenziale di trasportare molta più informazione di un cavo coassiale. Velocità di trasmissione superiori a 2 G bit/s su decine di chilometri sono state dimostrate, in confronto ad un massimo di 50 M bit/s su un chilometro per un cavo convenzionale.

Dimensioni inferiori e peso minore: le fibre ottiche sono considerevolmente più piccole dei soliti cavi. Il vantaggio qui è ovvio, specialmente laddove lo spazio è prezioso, come nei condotti telefonici sovraffollati che corrono sotto le strade delle città. Una fibra di 0,005 pollici di dia-

metro con un rivestimento di circa 0,25 pollici può, per esempio, trasportare la stessa informazione (e quindi sostituire) un mazzo di 900 paia di cavi di rame.

Il vetro pesa meno del rame: considerando le dimensioni più ridotte delle fibre, il risparmio in peso può essere significativo. I militari, per esempio, hanno trovato che un rimorchio da 1,25 tonnellate può trasportare fibre ottiche della stessa capacità di tre autocarri da 2,5 tonnellate carichi di cavo elettrico.

Minore attenuazione: le fibre presentano perdite di potenza del segnale più basse dei cavi coassiali. Inoltre, l'attenuazione nella fibra non cresce con la frequenza come avviene con i conduttori di rame. Sono state dimostrate attenuazioni di meno di 0,4 dB/km.

Insensibilità a interferenze elettromagnetiche (EMI) e a impulsi elettromagnetici (EMP): poiché è un materiale dielettrico, la fibra non è influenzata dai campi elettromagnetici e dai problemi collegati. Questo elimina la necessità di ingombranti e costose protezioni. Un alto grado di immunità allo EMP è un altro specifico vantaggio della fibra ottica rispetto ad un conduttore metallico.

Sicurezza: poiché le sue proprietà dielettriche offrono isolamento elettrico, la fibra non è soggetta a cortocircuiti, scintille e relativi pericoli. Non sono necessari dispositivi di protezione contro i fulmini per un cavo di fibre ottiche come è richie-

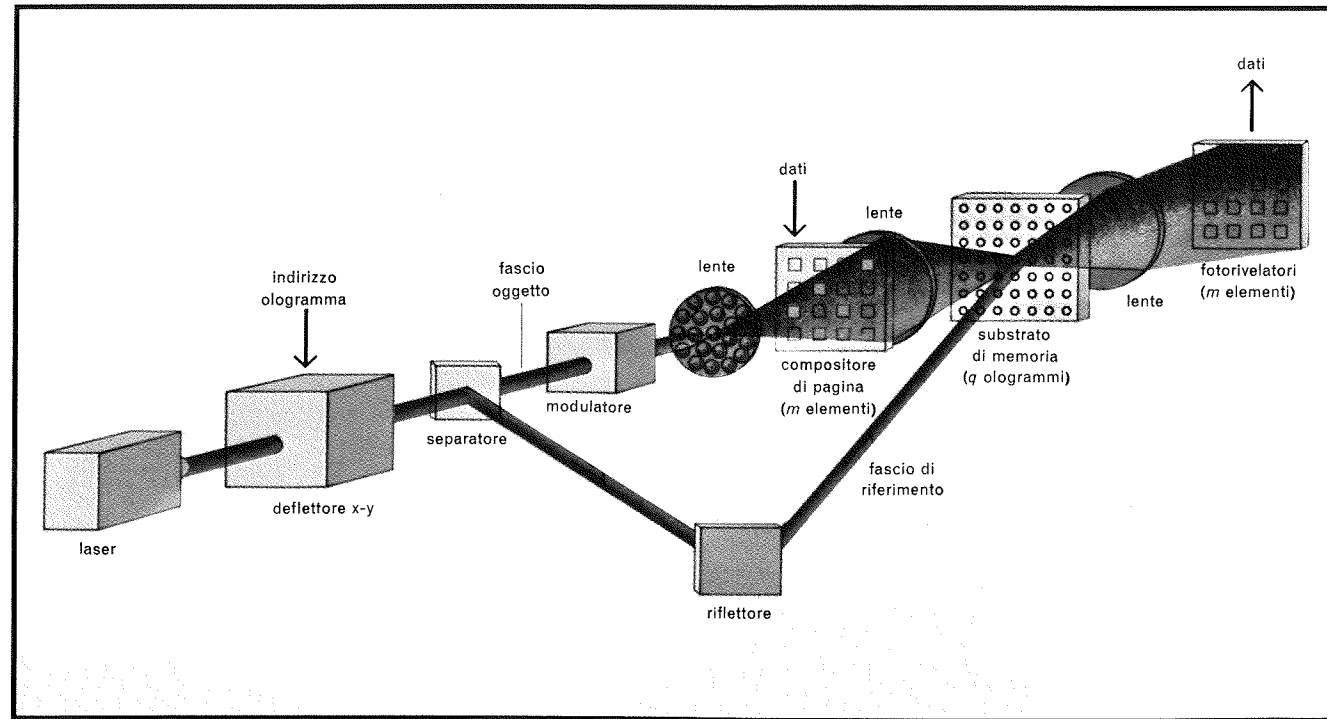


Fig. 6 - La figura illustra lo schema di una memoria ottica olografica, con substrato cancellabile. Sistemi di questo tipo avevano suscitato molte aspettative negli anni '70, che si sono però rivelate decisamente premature.

Riteniamo tuttavia che gli esempi fatti siano sufficienti per dare a tutti la sensazione della ricchezza di idee, della fantasia e dell'inventiva che hanno caratterizzato lo sviluppo dell'elaboratore. E se gran parte delle proposte non si sono alla fine affermate, esse hanno comunque portato un prezioso contributo di conoscenze e di esperienza. Il progresso tecnologico si vale anche degli "insuccessi".

Nota:
Questo articolo è apparso su ZERO UNO, sett. 1983, e viene qui riprodotto per gentile concessione dell'editore Mondadori.

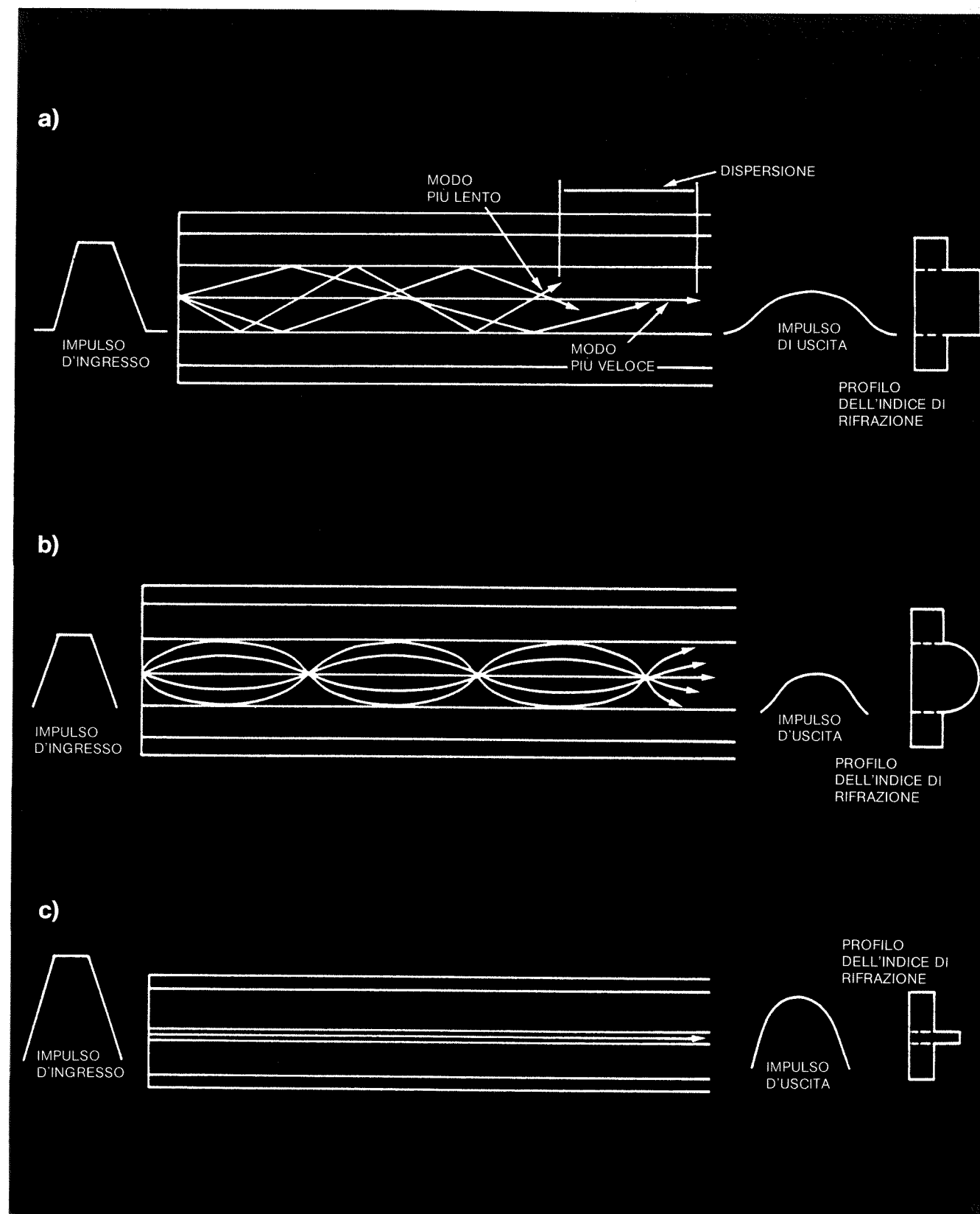


Fig. 1 - I vari tipi di fibre ottiche sono caratterizzate dal diverso profilo dell'indice di rifrazione: a) fibra multimode step-index, b) fibra multimode graded-index, c) fibra single-mode step-index. I tre tipi di fibre consentono distanze e velocità operative crescenti da a) a c), ma anche le difficoltà di impiego e i fattori di costo aumentano nello stesso ordine.

sto per il cavo metallico. In molti casi le fibre non devono essere alloggiare in condutture per rispettare degli standard restrittivi che si applicano ai cavi metallici e questo significa minori spese.

Disturbi/riservatezza: le fibre non irradiano energia e perciò non generano interferenze che possono disturbare altre parti del sistema. Inoltre, la necessità di protezione è eliminata e non è richiesto di soddisfare a onerose norme FCC. Oltre a ciò, la fibra non si presta intrinsecamente a intercettazioni e questo la rende un supporto sicuro, che può eliminare la necessità di costosi dispositivi di cifratura. Una fibra interrotta non crea un ambiente rumoroso per dispositivi elettronici, come avviene per un filo. Un sistema di fibre ottiche non soffre di problemi di interferenza o di ritorni di terra come avviene per un sistema basato sul metallo. E tutto ciò semplifica il progetto e riduce i costi del sistema.

Costo inferiore: benché i suoi componenti possano essere oggi più costosi, il collegamento via fibra ottica può essere alla fine meno costoso di un sistema metallico. La dimensione e il peso inferiori significano costi inferiori di gestione e di immagazzinamento. L'isolamento elettrico riduce i costi di schermatura. La bassa attenuazione riduce il numero di ripetitori necessari per rigenerare il segnale su lunghe distanze. In aggiunta, riduzioni di costo significative si verificano man mano che la tecnologia matura; il vetro e la plastica, dopo tutto, sono più economici del rame. Inoltre, i vantaggi delle fibre ottiche spesso giustificano costi aggiuntivi.

2. La fibra

Le caratteristiche di trasmissione ottica delle guide d'onda in fibra ottica sono espresse in termini di attenuazione o perdita e di larghezza di banda o dispersione dell'impulso.

In fibre fatte di vetro al silicio, i meccanismi che causano perdite sono l'assorbimento e la dispersione. La dispersione di un tratto di fibra è governata dalla differenza nel tempo di

arrivo dell'impulso di energia ottica spedito ad una estremità e ricevuto all'altra. Il progetto della fibra può alterare sostanzialmente la dispersione. La dispersione è causata dalle caratteristiche di propagazione della guida d'onda, dalla larghezza spettrale della sorgente di luce e dalla forma della fibra. In contrasto con i cavi di rame in cui l'attenuazione aumenta con la radice quadrata della larghezza di banda, la fibra offre attenuazione costante per ogni larghezza di banda operativa.

Per introdurre quanta più luce è possibile dal diodo entro la fibra, l'area attiva del diodo dovrebbe avere un diametro non superiore a quello della fibra. Altrimenti buona parte della luce entra nel rivestimento esterno e non contribuisce utilmente alla potenza del segnale propagato. Così quanto più alta è l'apertura numerica, tanto più efficiente è l'accoppiamento. La penalità per il migliore accoppiamento è la dispersione causata dalla maggiore differenza di cammino tra i nodi estremi propagati.

L'attenuazione spettrale è un'altra importante considerazione. Molto semplicemente, alcune lunghezze d'onda viaggiano sulla fibra meglio di altre. Un'onda di 820 nm viaggia bene lungo la fibra, mentre un'onda di 950 nm è fortemente attenuata. L'attenuazione in una fibra va indicata per una certa lunghezza d'onda; per es. 5 dB/km a 820 nm.

L'attenuazione spettrale della fibra solleva due importanti argomenti: l'accoppiamento dei componenti e la purezza dello spettro. Se una fibra propaga bene un'onda di 820 nm, la sorgente deve fornire una forte energia luminosa a questa lunghezza d'onda e il rivelatore dovrebbe essere molto sensibile a questa stessa lunghezza d'onda. Attualmente sono disponibili componenti per l'intervallo attorno a 820 nm. Le fibre presentano perdite ancora minori nella regione infrarossa attorno a 1200 e 1500 nm; sorgenti e rivelatori per questo intervallo sono al momento poco disponibili, ma comunque in corso di sviluppo.

3. Sorgenti

In molti sistemi di comunicazione ottica, sorgenti di luce a semiconduttori vengono usate per convertire i segnali elettrici in luce introdotta poi nelle fibre.

Quando la giunzione p-n di un diodo a semiconduttore è opportunamente polarizzata, ha luogo un'emissione spontanea di fotoni. La lunghezza d'onda del fotone è determinata dall'intervallo di energia tra le bande di valenza e di conduzione nel semiconduttore. Nel Ga Al As questo corrisponde a circa 850 nm e può essere variata lievemente, modificando la percentuale di alluminio.

I laser permettono di inviare nel nucleo della fibra la massima quantità di potenza ottica. Un laser a semiconduttori è formato confinando i fotoni emessi nella regione p-n mediante specchi realizzati sulle estremità del dispositivo. L'azione di amplificazione del laser si mantiene una volta superata la corrente di soglia necessaria per fornire una densità elettronica sufficientemente elevata. I LED sono preferibili ai laser poiché hanno migliore affidabilità, costano di meno e non richiedono compensazione di temperatura, come i laser.

4. Rivelatori

Il rivelatore svolge una funzione opposta a quella della sorgente. Esso converte il segnale luminoso in un segnale elettrico. Il circuito di uscita del ricevitore amplifica il segnale e riproduce accuratamente il segnale digitale originale. Normalmente, un ricevitore limita la velocità di un collegamento in fibra ottica.

Ci sono quattro tipi fondamentali di rivelatori: il fotodiodo PIN, il fotodiodo a valanga (APD), il fototransistore e il transistor fotodarlington. Ciascuno ha i propri vantaggi e svantaggi. Sensibilità, velocità e costo sono probabilmente i tre elementi principali da considerare. Attualmente la maggioranza dei sistemi a fibre ottiche usa dei fotodiodi PIN. Se si richiede sensibilità e velocità maggiori, si possono usare i più costosi APD.

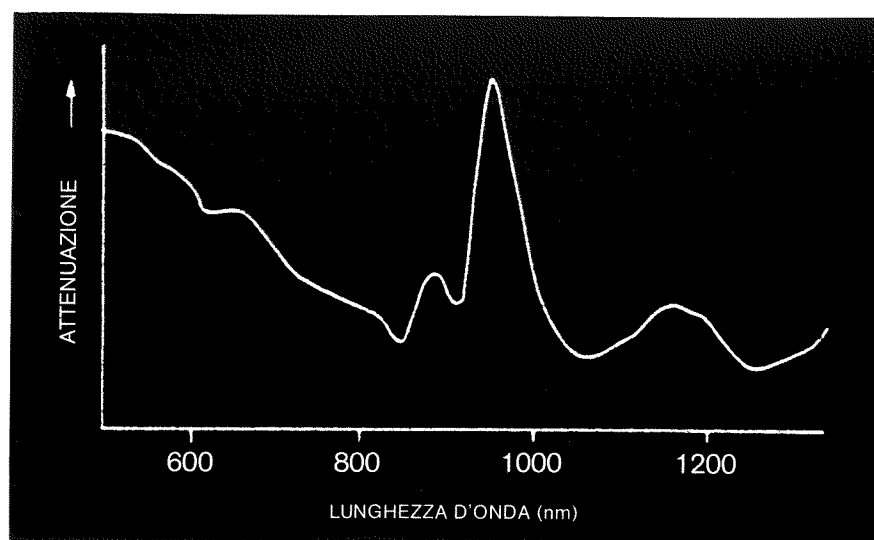


Fig. 2 - Attenuazione di una fibra ottica in funzione della lunghezza d'onda della luce.

4. Sistemi

Per un sistema a fibre ottiche, l'architettura più semplice è una configurazione punto a punto. Per comunicazioni bidirezionali si può usare un cavo o due fibre. Oppure, si possono usare accoppiatori ottici direzionali con una sola fibra.

Sfortunatamente, questa semplice architettura non si adatta ad alcune applicazioni. Molto spesso, i sistemi richiedono comunicazione tra più terminali geograficamente distribuiti. In questi casi occorre una configurazione ad accesso multiplo (spesso chiamata una struttura a bus o un sistema a molti terminali). Questa architettura richiede un componente addizionale cioè, un accoppiatore. Ci sono due possibili accoppiatori, quello a T e quello a stella.

Con un generico accoppiatore a T, i singoli utenti possono prelevare o iniettare energia ottica entro la linea principale in corrispondenza di ogni accoppiatore. La configurazione a stella offre ai progettisti una selezione alternativa di accoppiamento per applicazioni con molti terminali. La Optoelectronics Division della Honeywell produce accoppiatori a stella fino a 64 porte.

Con la configurazione in-linea, il segnale ottico subisce una perdita ad ogni accoppiatore a T del bus. La perdita nel caso pessimo cresce linearmente con il numero dei terminali. D'altro lato, la configurazione a stella introduce soltanto la perdita di un accoppiatore. Inoltre, al crescere

del numero n delle porte la perdita dovuta all'accoppiatore a stella cresce soltanto come $10 \log(n)$. Usando delle unità con una perdita costante di prelievo, gli accoppiatori a T possono anche creare dei problemi dinamici. Si è dimostrato che il dispositivo a stella ha perdite minori per sistemi con più di quattro terminali, all'incirca.

5. Multiplexing

La trasmissione di più segnali individuali richiede qualche forma di multiplexing. Una tecnica possibile è il multiplexing a divisione di spazio che usa una fibra separata per ogni segnale. Tuttavia il costo delle fibre diventa presto proibitivo se si devono trasmettere molti segnali. Il multiplexing di frequenza e a divisione di tempo può invece essere usato per convogliare più canali entro una sola fibra.

Una tecnica relativamente nuova, il multiplexing a divisione di lunghezza d'onda (wavelength division multiplexing: WDM) usa sorgenti di lunghezza d'onda distinta per separare i vari canali. Benché si richieda un filtraggio ottico ai fotorivelatori per selezionare il canale appropriato, il sistema WDM accresce ulteriormente il vantaggio di larghezza di banda delle fibre ottiche. È uno schema passivo di multiplexing, che facilita la crescita e permette operazioni con frequenza dei dati più bassa, che può dar luogo a costi minori del sistema.

6. Prospettive future

Durante il decennio in corso ci sarà un ulteriore progresso nella tecnologia delle fibre ottiche. Selettori, dispositivi di commutazione ottica e configurazioni integrate verranno sviluppati parallelamente. Questi miglioramenti ridurranno ulteriormente il costo e aumenteranno i vantaggi tecnici dei sistemi a fibre ottiche, il che, a sua volta, ne estenderà il campo di applicazione.

Il ruolo più importante delle fibre ottiche sarà quasi certamente nelle reti di telecomunicazioni. Questa previsione è basata sull'idea generalmente accettata che la rete del futuro sarà completamente digitale: sarà cioè una "Integrated Services Digital Network" (ISDN) in cui sia i servizi a banda ristretta che quelli a larga banda verranno distribuiti ai singoli abbonati con un unico supporto di trasmissione. Senza dubbio la fibra ottica sarà il mezzo di trasmissione prescelto.

I costi delle fibre ottiche e dei relativi componenti stanno calando rapidamente. Entro alcuni anni il collegamento in rete locale (LAN) sarà economicamente fattibile, anche se inizialmente verranno offerti soltanto i servizi a banda ristretta. Sembra esserci ormai una tendenza significativa di passaggio dalle reti LAN metalliche a reti configurate con fibre ottiche per permettere una futura espansione del sistema e l'incorporazione di servizi aggiuntivi.

L'attività sulle reti LAN a fibre ottiche si è finora concentrata sul come

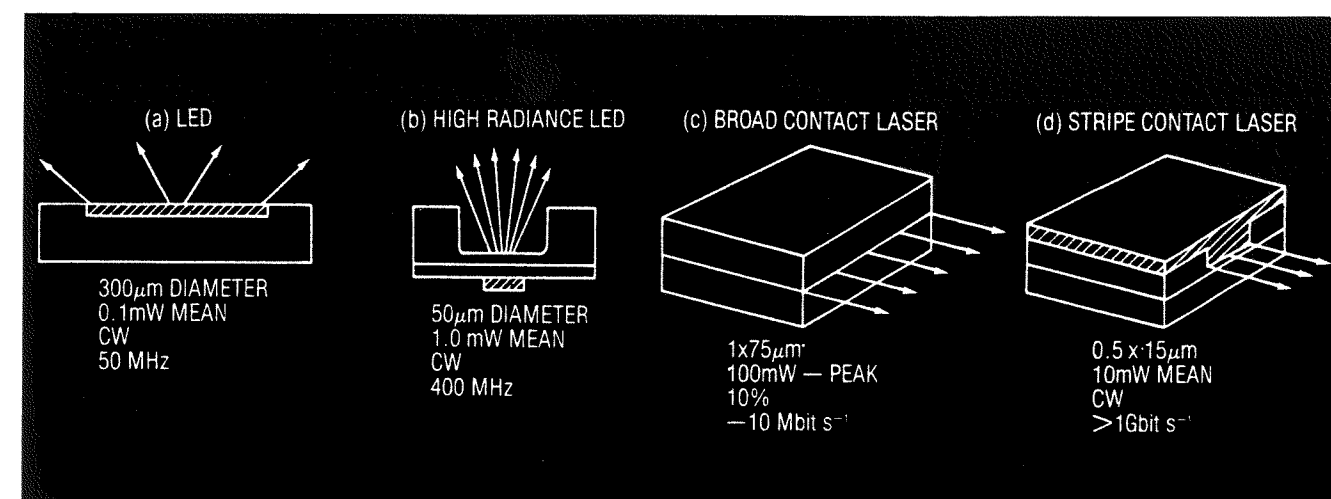


Fig. 3 - Vari tipi di sorgenti di luce, a LED e a LASER.

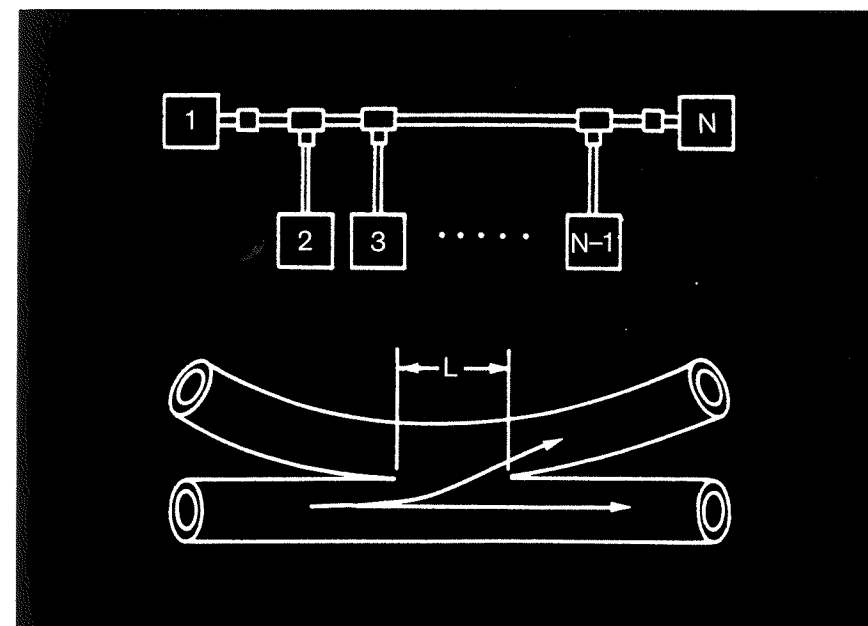


Fig. 4 - Accoppiatori ottici a T. La figura in alto mostra un generico accoppiatore a T. Quella in basso, un accoppiatore a T in cui le fibre sono fuse insieme; in questo caso, la quantità di luce trasferita varia in funzione della lunghezza di accoppiamento L .

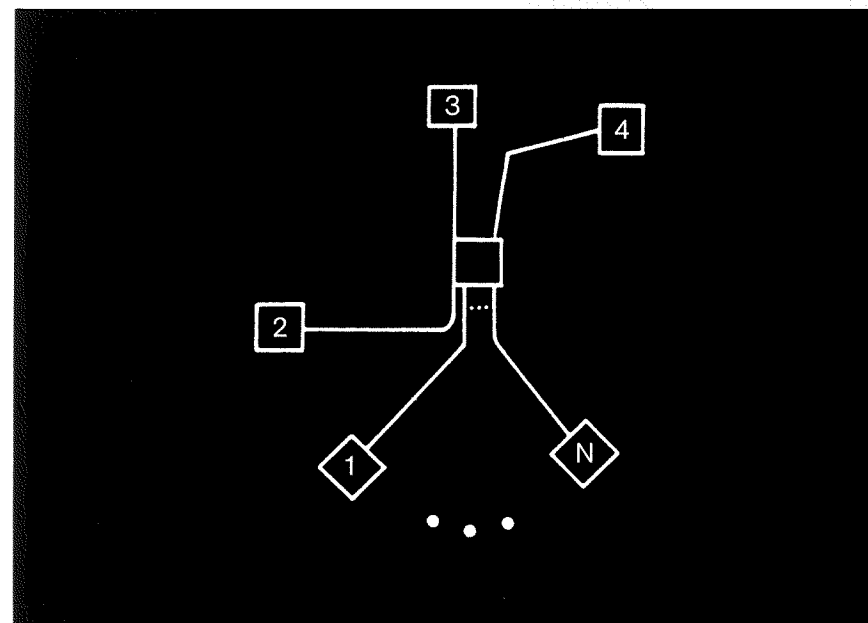


Fig. 5 - Accoppiamento a stella. Questa soluzione presenta perdite minori della precedente per un numero di terminali superiore a quattro.

adattare le reti a fibre ottiche agli standard dei protocolli orientati alle reti metalliche. Il passo successivo è il progetto di reti LAN basate sul trarre il massimo beneficio dalle fibre ottiche, minimizzando le limitazioni. Per esempio, le larghezze di banda estremamente ampie ottenibili dai componenti in fibre ottiche comunemente a disposizione possono essere usati per realizzare protocolli di accesso più semplici, anche se inefficienti per la larghezza di banda, riducendo così i costi di interfaccia. Un altro esempio è nell'area dei calcolatori. Qui l'uso efficiente della trasmissione dati via fibre ottiche per migliorare la velocità dell'interconnessione dei calcolatori richiede modifiche fondamentali negli attuali protocolli paralleli di input/output. Tuttavia, una graduale evoluzione verso nuovi protocolli è essenziale all'accettazione di queste tecniche.

7. Conclusioni

Chiaramente, i sistemi di comunicazione a fibra ottica hanno dimostrato sul campo le loro capacità e stanno guadagnando ampia accettazione. La loro convenienza economica è sicura per sistemi a velocità di trasmissione relativamente alta e si è fiduciosi che essi offriranno l'approccio più economico, flessibile ed espandibile, nella maggioranza delle applicazioni per linee principali e per lunghe distanze. Inoltre, crescendo la domanda per servizi a larga banda, i sistemi a fibre ottiche saranno la scelta logica, o anche l'unica possibile, per circuiti ad anello e per reti locali. La comunicazione via fibre ottiche giocherà pertanto un ruolo sempre maggiore nell'evoluzione verso le reti digitali dei servizi integrati del futuro.

8. Bibliografia

- [1] Hawkins, Bobby; "The Development of Fiber Optic Components", *Scientific Honeywell*, Vol. 1, No. 3, (September, 1980), pp. 2-11.
- [2] Kao, Charles K.; *Optical Fiber Systems: Technology, Design, and Applications*, McGraw-Hill Book Company, 1982.
- [3] Sandbank, C.P.; *Optical Fiber Communication Systems*, John Wiley and Sons, 1980.
- [4] Husain, A., Cook, S.D.; *Impact of Fiber Optics on Local Area Networks*, SPIE San Diego, August 1982.

NOTA. Il materiale qui presentato è stato tratto dall'articolo "Fiber Optic Communications" dello stesso Autore, pubblicato su "SCIENTIFIC HONEYWELL", vol. 3, n. 4, Dec. 1982.

Tecniche di individuazione degli errori nella trasmissione dati

TREVOR HOUSLEY

A tutti noi è capitato di sostenere conversazioni telefoniche disturbate, anche gravemente, da rumori di fondo, ticchettii e brusii o addirittura interferenze durante le quali abbiamo potuto ascoltare le conversazioni altrui; tutti questi disturbi non desiderati sono *rumore*.

Il cervello umano è un ottimo calcolatore e si adatta alle mutevoli e varie condizioni della linea e riesce a minimizzare l'effetto del rumore. Per esempio se un gracchio sulla linea cancella una parola, di solito siamo in grado di sostituire a quel rumore la parola perduta senza essere sempre costretti a chiedere all'altra persona di ripetere. Se tuttavia le condizioni della linea peggiorano, potremmo chiedere all'interlocutore di parlare più lentamente e magari più chiaramente. Con tutte queste azioni, mittente e destinatario cercano un modo per adattarsi alle mutevoli condizioni della linea telefonica.

I rumori che si odono sulla linea telefonica vengono causati dalle commutazioni che avvengono nelle centrali telefoniche e da cause esterne, elettriche o magnetiche, come fulmini, cadute di potenza, lavori di riparazione sulla linea, o dipendenti da altre linee e da altri impianti. Gli impulsi di rumore raccolti dalle linee telefoniche vengono percepiti generalmente come scariche che di solito si distribuiscono in modo casuale. Gli errori invece si distribuiscono in modo sistematico, cioè ad intervalli regolari di tempo, o secon-

do qualche altro schema che di solito è prevedibile, come un malfunzionamento dello hardware che nelle stesse circostanze ripete sempre lo stesso errore. E poiché è più facile prevederli, è anche più facile eliminare gli errori sistematici che non quelli casuali.

Poiché la maggior parte dei sistemi di trasmissione dati ricorre, come mezzo di trasmissione, alle linee telefoniche, questi sono soggetti allo stesso tipo di rumore che disturba le nostre conversazioni telefoniche. Purtroppo, però, gli elaboratori che gestiscono i sistemi di trasmissione dati non sono così brillanti come il cervello umano e se quindi perdono, nel corso della trasmissione di un flusso grezzo di dati, uno o più bit, non sono in grado di ricostruirli e neppure di scoprire gli errori, a meno che non abbiano delle informazioni ridondanti che consentano loro d'individuare errori od omissioni.

È necessario, nella maggior parte dei sistemi commerciali, individuare e correggere gli errori. In questo articolo esamineremo i mezzi a nostra disposizione per scoprire gli errori; nell'ultimo paragrafo inizieremo l'analisi delle modalità di correzione degli errori individuati.

In molti sistemi con controllo e verifica continua, è sufficiente scoprire che si è verificato un errore senza doverlo poi correggere. I sistemi d'inseguimento dei satelliti in orbita

attorno alla Terra, schematizzati nella fig. 1, sono di questo tipo: il sistema radar tiene sotto controllo la posizione del satellite rilevandone diverse posizioni durante ogni secondo. I dati di ogni rilevazione sono poi trasmessi lungo una linea di comunicazione all'elaboratore centrale. Se durante quel percorso un blocco di dati incontra un disturbo o rumore, esso verrà rovinato e conterrà degli errori, ma è sufficiente che il sistema destinatario scopra l'errore e rifiuti di conseguenza i dati errati. Non è necessario correggere l'errore poiché l'orbita del satellite è relativamente stabile e il vero valore del punto che avrebbe dovuto essere rilevato e che non è giunto a causa dell'errore può essere calcolato dall'elaboratore destinatario interpolando e proiettando i dati ricevuti correttamente sino a quel momento e relativi alle posizioni occupate in precedenza dal satellite stesso. Si tratta di una situazione generalmente vera per tutti i sistemi di controllo continuo.

Se invece i dati trasmessi e giunti errati fossero stati dati amministrativi, come gli elementi per calcolare paghe e stipendi, avremmo dovuto scoprire il o i dati errati (non sarebbe possibile infatti correggere quei dati interpolando semplicemente i dati immediatamente precedenti o successivi perché in un sistema amministrativo non è detto che esista alcuna relazione tra i dati contenuti in un blocco e quelli del blocco successivo).

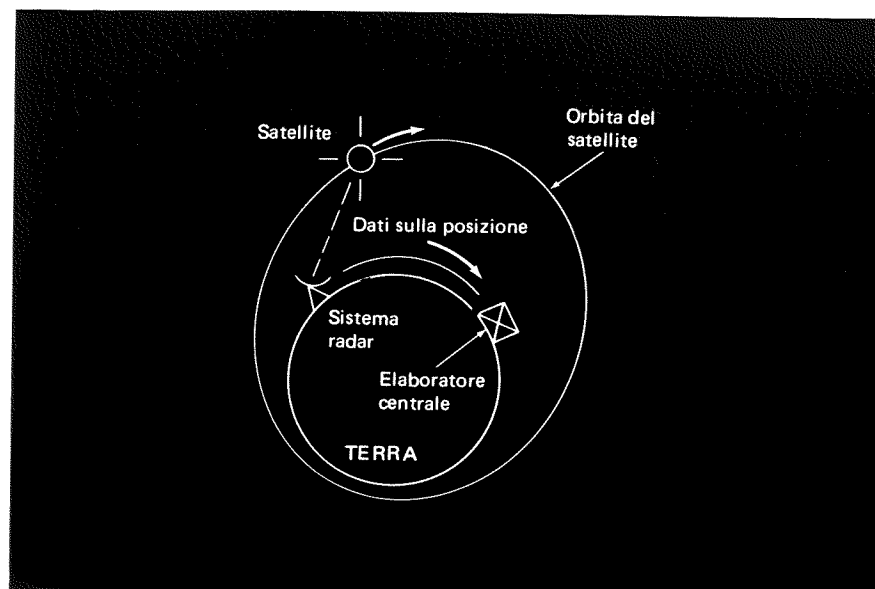


Fig. 1 - Sistema di inseguimento dei satelliti

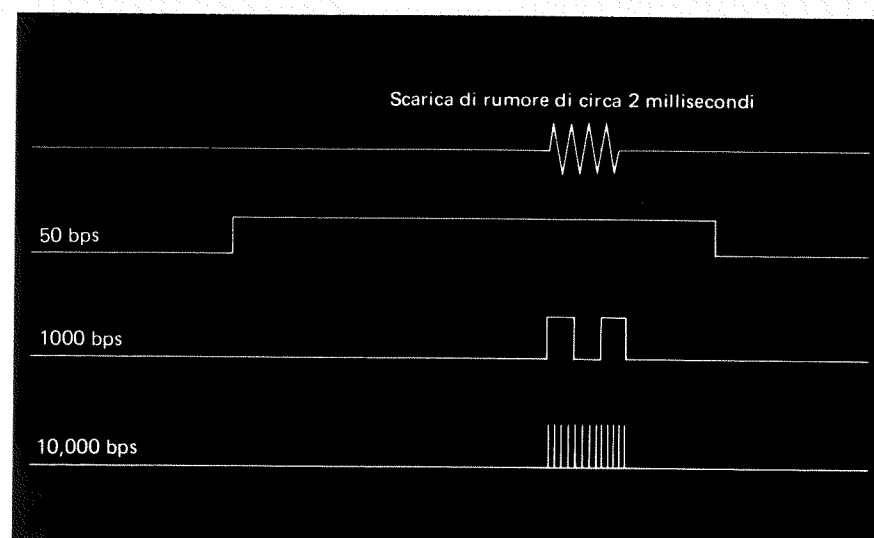


Fig. 2 - L'effetto del rumore sui dati a diverse velocità di trasmissione

1. Frequenza degli errori

La frequenza degli errori sui sistemi di trasmissione dei dati varia al variare della velocità di trasmissione. La fig. 2 indica l'effetto che può avere una scarica di rumore di 2 millisecondi sui dati che vengono trasmessi a diverse velocità. Alla velocità di 50 bit al secondo, ogni bit ha una durata di 20 millisecondi. L'hardware in ricezione campiona ogni bit il più possibile vicino al centro del bit stesso per determinare se è un 1 o uno 0: per determinare lo stato del bit è sufficiente un istante. Quindi se la scarica di rumore altera il bit mentre il destinatario sta campionando lo stato del bit, questo può percepire o meno l'errore a seconda del momento preciso in cui avviene la sca-

rica e alla velocità di 50 bit al secondo è assai improbabile che una scarica di rumore della durata di 2 millisecondi possa alterare i dati, perché c'è solo una probabilità su dieci di percepire la scarica di rumore proprio nell'istante in cui si campiona il bit.

Se però si aumenta la velocità di trasmissione a 1000 bit al secondo, ogni bit dura un millisecondo e la scarica della durata di due millisecondi interesserà due bit ed è assai probabile che uno o due bit successivi siano distorti. Aumentando la velocità di trasmissione a 10.000 bit al secondo, la scarica di due millisecondi durerà un tempo pari a quello durante il quale vengono trasmessi 20 bit e si è quasi certi che uno o più bit verranno alterati.

È evidente quindi che le probabilità di errore aumentano all'aumentare della velocità di trasmissione. Inoltre i modem ad alta velocità di solito impiegano tecniche complesse di modulazione per cui è possibile che un solo impulso di rumore alteri una stringa di bit o che possa addirittura interrompere la demodulazione.

Normalmente è difficile farsi dire dalle società telefoniche quali sono le frequenze con cui si hanno errori sulle linee perché queste non si erano particolarmente preoccupate degli impulsi istantanei di rumore fino a quando non si è cominciato ad usare le reti telefoniche per la trasmissione dei dati. La tab. 1. ci dà un'indicazione della frequenza di errore che ci si può aspettare sce-

gliendo a caso una linea da una normale rete telefonica e trasmettendo a caso dati lungo di essa a diverse velocità. Per la trasmissione a 9600 bps si è specificato un intervallo, anziché un unico valore, perché le reti telefoniche sono composite e sono costruite impiegando circuiti assai vari: coppie aperte di fili su pali, cavi sepolti sotto le strade delle nostre città e costituiti da coppie di fili intrecciati, cavi coassiali ad alta capacità e qualità, collegamenti via microonde radio o tramite satelliti. A seconda che, al momento della chiamata telefonica, ci capitino componenti di buona o di cattiva qualità avremo frequenze di errore più vicine all'estremo migliore o a quello peggiore dell'intervallo indicato. A velocità di trasmissione inferiori, i sistemi sono relativamente insensibili alle variazioni della qualità del cavo.

Prove eseguite negli Stati Uniti (*Bell System Technical Journal*, aprile 1971) hanno indicato che le frequenze di errore variano enormemente da una linea all'altra - anche di tre o quattro ordini di grandezza. Pertanto occorre interpretare le cifre della tab. 1 con molta prudenza: esse rappresentano i valori che di solito speriamo d'ottenere con una rete telefonica progettata secondo le raccomandazioni Ccitt.

La tabella dimostra anche perché in molti paesi non è possibile trasmettere dati a velocità superiore ai 1200 bps sulle reti telefoniche commutate: infatti, se la frequenza di errore è superiore a uno su centomila, la rete telefonica deve essere considerata in genere inadatta alla trasmissione dei dati. Eppure si mantengono molte reti su linee con una frequenza di errore anche superiore: questa apparente contraddizione si spiega col fatto che gli errori tendono a manifestarsi come scariche (cioè i bit errati non tendono a distribuirsi a caso: sono invece le scariche di errore che si distribuiscono in modo casuale). Quando si trasmettono i dati in blocchi, si trova che ciò consente di trarre qualche vantaggio, perché un blocco con un bit errato di solito non può essere utilizzato così come del resto avviene per i blocchi con 10 o 15 bit errati. Pertanto se, per esem-

pio, si verificassero 100 impulsi di rumore della lunghezza di un bit, distribuiti casualmente in un dato periodo, si potrebbe pensare ragionevolmente di trovare 100 blocchi inutilizzabili. Se però lo stesso numero di impulsi di rumore fosse distribuito in gruppi o scariche sullo stesso periodo, il numero di blocchi inutilizzabili sarebbe significativamente inferiore.

Si può poi verificare facilmente che in presenza di elevate frequenze di errore è più conveniente trasmettere per blocchi di ampiezza limitata che hanno maggior probabilità di superare le scariche di errore che non i blocchi di dati più lunghi.

Tab. 1 - Frequenze statistiche tipiche degli errori su una linea telefonica scelta a caso.

Velocità di trasmissione	Frequenza degli errori casuali in bit
1200 bps	1 su 200.000
2400 bps	1 su 100.000
9600 bps	da 1 su 1000 a 1 su 10000

2. Condizionamento della linea

Al momento della locazione di una linea telefonica che dovrà essere collegata in modo permanente a terminali ed elaboratore, è possibile chiedere e ottenere dalla società telefonica di costruire speciali «ponticelli» che consentano di evitare le stazioni di commutazione delle centrali telefoniche, riducendo così la quantità di rumore sulla linea. La società telefonica può anche misurare le prestazioni della linea prescelta e, con opportuni dispositivi elettrici, migliorarne le caratteristiche. Questo processo, chiamato *condizionamento*, permette d'aumentare le caratteristiche della linea riducendo le frequenze con cui si verificano scariche che danno luogo ad errori, consentendo di impiegare le linee per velocità elevate - ad esempio 9600 bps - con frequenze di errore accettabili. Si possono condizionare solo le linee private poiché non è possibile condizionare le linee telefoniche commutate dato che la selezione delle tratte fisiche che vengono unite dalle centrali telefoniche per costruire il collegamento tra chiamante e chiamato su rete telefonica com-

mutata non è nota a priori e avviene in modo apparentemente casuale.

3. Equalizzazione

I modem moderni sono in grado di controllare le condizioni della linea e dei segnali in arrivo e reagire automaticamente ad alcuni scostamenti del comportamento normale, cioè svolgono, spesso in modo automatico, le funzioni chiamate di *equalizzazione*. Oggi però non si è ancora in grado d'ottenere gli stessi risultati che si possono avere con il condizionamento della linea, anche se così è possibile trasmettere in pratica a 4800 bps impiegando linee di normali reti commutate.

Spesso sulle multidrop si ricorre congiuntamente all'equalizzazione e al condizionamento. Per il modem centrale (*instation modem*) di una linea multidrop, le caratteristiche della linea variano al variare del modem esterno o periferico (*outstation modem*) in trasmissione. In questo caso il condizionamento della linea consente di raggiungere un buono standard qualitativo, mentre l'equalizzazione adattativa del modem centrale permette di compensare le diverse caratteristiche delle linee periferiche.

4. L'individuazione degli errori

Il modo più semplice per gestire gli errori è il non individuarli del tutto. Di solito, come abbiamo detto, ciò non è accettabile, ma esistono sistemi in cui i soli dati trattati sono costituiti da testi scritti e può risultare economico e facile lasciare all'operatore umano il compito di leggere, interpretare e correggere il messaggio inviato: questi, infatti, può, seguendo un procedimento complesso, manipolare il linguaggio riuscendo anche a interpretare una frase confusa e a renderla ordinata aiutandosi con la sintassi e il contesto generale del messaggio. Per esempio la frase:

Si prevede un peggioramento delle condizioni del tempo

può essere correttamente interpretata come:

Si prevede un peggioramento delle condizioni del tempo

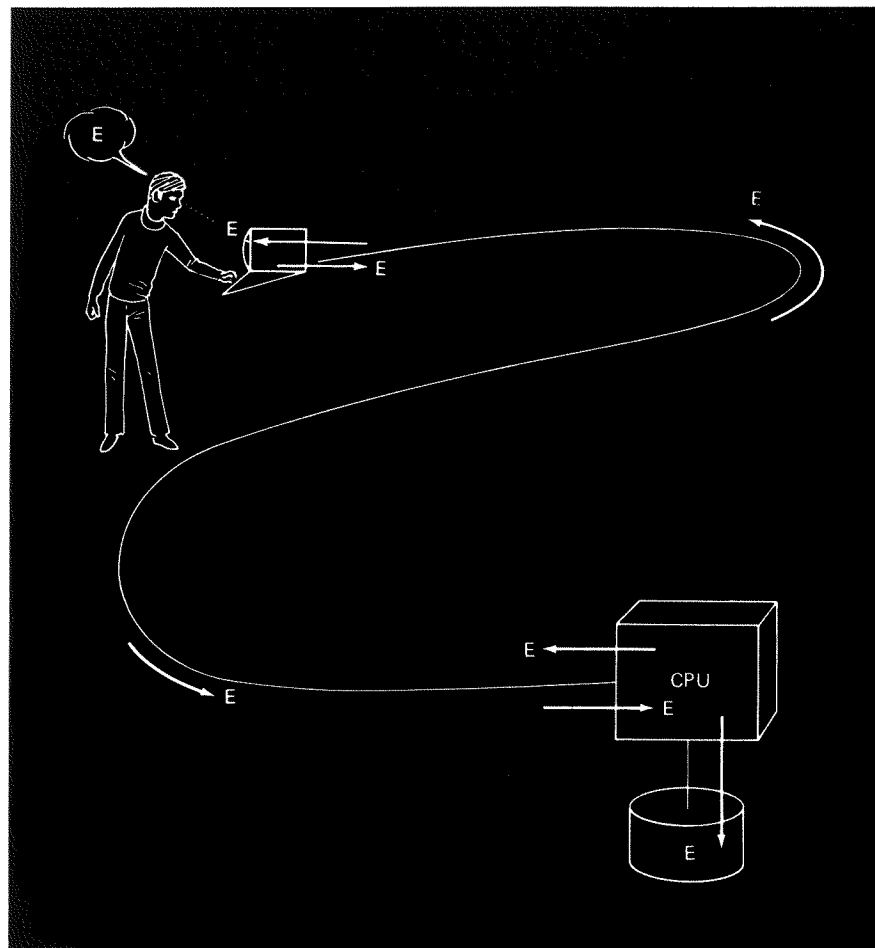


Fig. 3 - La tecnica dell'eco

Se non ci sono troppi errori, la maggior parte delle persone è in grado di solito d'interpretare correttamente il testo del messaggio. I sistemi telegrafici e telex in genere si avvalgono di questa possibilità, ed è per questo che normalmente si ripetono alla fine del telegramma le informazioni significative come i numeri perché è più difficile ricostruire intuitivamente le informazioni numeriche. Per individuare gli errori nei sistemi di trasmissione dei dati è necessario ricorrere a procedure o mezzi ridondanti (cioè inviare altre informazioni oltre a quelle strettamente necessarie per trasmettere il testo): tanto maggiore sarà il livello di ridondanza, tanto più affidabile sarà il metodo di individuazione dell'errore. Ma poiché l'informazione ridondante impiega parte della capacità trasmissiva che potrebbe essere altrimenti utilizzata per trasferire i dati, la maggior parte dei sistemi d'individuazione degli errori è progettata in modo da raggiungere un compromesso accettabile tra costo delle informazioni ridondanti che sarebbero ri-

chieste per elevare ulteriormente la qualità della trasmissione e il beneficio che si potrebbe trarre aumentando ulteriormente la percentuale di errori individuati.

1. La tecnica dell'eco

La *tecnica dell'eco* (chiamata talvolta *ecoplex* o *echo technique*) è un metodo assai semplice per individuare gli errori a cui si ricorre spesso nei sistemi interattivi (che sono quelli in cui l'operatore invia informazioni all'elaboratore per mezzo di telescriventi o terminali video), come è schematizzato nella fig. 3.

In questo caso l'operatore pensa al carattere che vuole inviare all'elaboratore e batte il tasto corrispondente sul terminale. Il terminale, a sua volta, trasmette il carattere all'elaboratore che lo riceve e lo memorizza su disco e quindi lo ritrasmette - creando l'eco - al terminale sulla stessa linea di trasmissione dati. Il carattere verrà quindi ricevuto dal terminale che lo visualizza consentendo

all'operatore di verificare se è proprio il carattere che voleva inviare. Se si è verificato un errore di comunicazione lungo il percorso, è probabile che il carattere sia stato distorto: in tal caso l'operatore troverà che il carattere rinviato come eco e quindi visualizzato sullo schermo non corrisponderà a ciò che voleva trasmettere e potrà così individuare e correggere l'errore.

In alcune rare occasioni può capitare che il carattere inviato sia stato distorto da un rumore e durante il percorso di eco sia stato ritrasformato in quello stesso carattere pensato e voluto dall'operatore. In tal caso non si riuscirà ad individuare l'errore: da un lato l'elaboratore memorizzerà il carattere errato ricevuto come corretto e dall'altro invece, dopo la verifica, l'operatore riterrà che la trasmissione sia avvenuta correttamente, ma si tratta di un evento statisticamente assai improbabile e, in concreto, nella maggior parte dei casi di sistemi interattivi, si accetterà l'eventualità di errori con probabilità così basse.

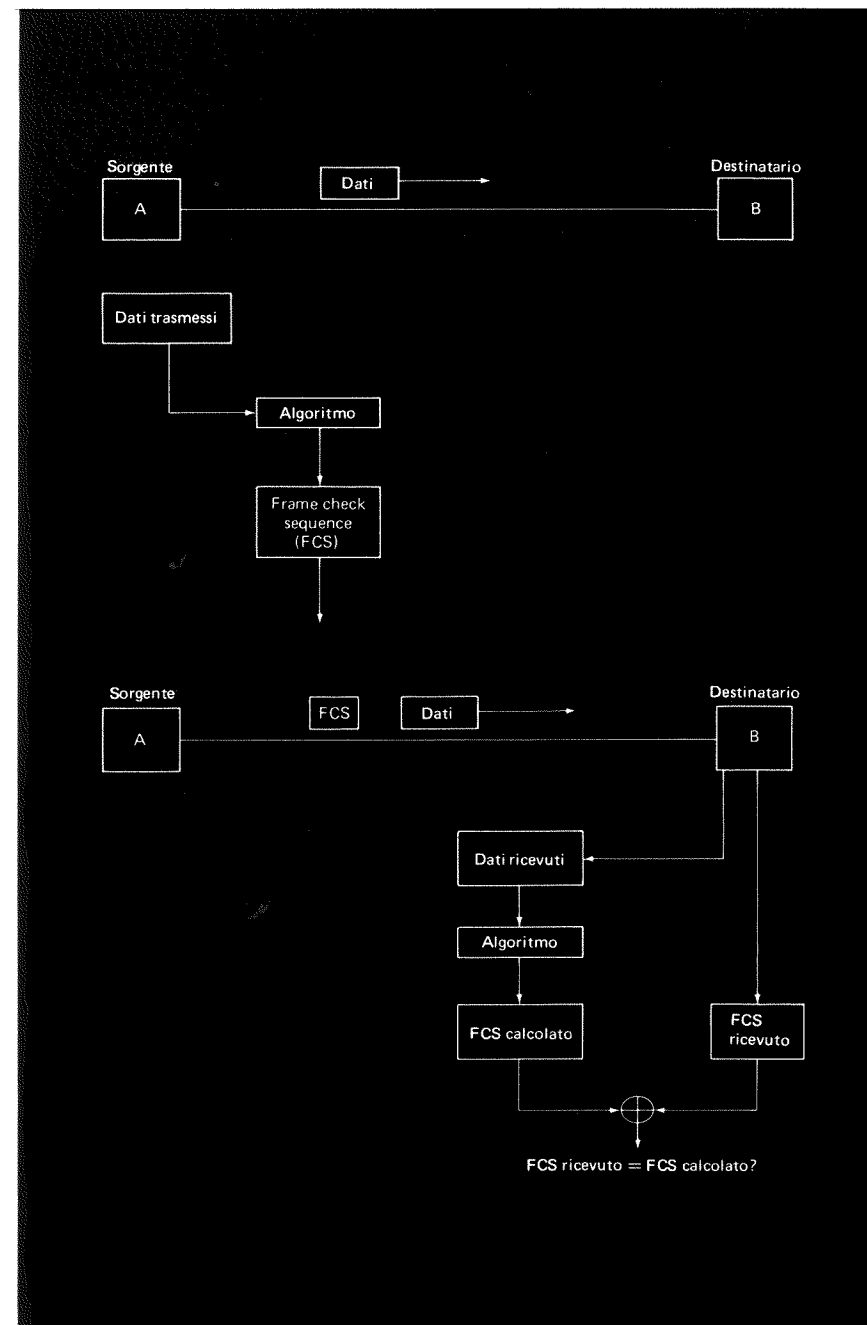


Fig. 4 - Metodo generale per l'individuazione automatica degli errori

5. Tecniche automatiche di individuazione degli errori

Sui sistemi di elaborazione dati più moderni si cerca di rendere il più automatico possibile l'individuazione e la correzione degli errori in modo da rendere minimo l'intervento dell'operatore e migliorare le prestazioni del sistema, eliminando i tempi di reazione relativamente lunghi tipici delle operazioni degli esseri umani.

L'individuazione automatica dell'errore si basa su un principio generale da cui derivano due metodi comunemente seguiti. La fig. 4 sche-

matizza il principio generale: i dati vengono trasformati matematicamente con un algoritmo, calcolando una *Frame Check Sequence* o FCS (sequenza di controllo del blocco): i dati, seguiti da FCS, vengono inviati lungo la linea. Il destinatario analizza i dati ricevuti mediante lo stesso algoritmo in modo da determinare il cosiddetto FCS calcolato, e quindi provvede a confrontare la FCS ricevuta dalla linea con quella calcolata; se corrispondono, i dati vengono considerati validi; in caso contrario il blocco di dati viene dichiarato errato e si dovrà intraprendere un'azione correttiva.

Il segreto del sistema sta tutto nel costruire un algoritmo tale da rendere assai improbabile che i dati errati o FCS superino il test. Vediamo ora i due metodi di impiego più frequenti: il *controllo di parità a due coordinate* e il *controllo ciclico di ridondanza*.

1. Controllo di parità a due coordinate

È noto che il normale controllo di parità sul carattere è da solo insufficiente, permettendo appena di verificare se il complemento al numero di bit contenuti in un carattere è pari o dispari. Quando l'elaboratore o un

terminale dotato di buffer trasmettono dati in blocchi, possiamo estendere le possibilità di verifica del semplice controllo di parità aggiungendo al blocco un *carattere di controllo del blocco* (BCC o *Block Check Character*) alla fine del blocco di dati. La fig. 5(a) rappresenta un blocco di dati in trasmissione, in cui ogni carattere ha un suo bit di parità: alla fine del blocco viene poi aggiunto il carattere di controllo del blocco che offre un controllo di parità cumulativa su tutti i caratteri precedenti, come si vede nella fig. 5(b). La fig. 5 riporta un blocco di 14 caratteri, per ognuno dei quali è indicato il bit di parità, e il carattere di controllo del blocco. Il bit 1 del carattere di controllo del blocco è il bit di controllo di parità calcolato su tutti i bit 1 dei caratteri del blocco. Il bit

4 del carattere di controllo del blocco è un bit di controllo di parità calcolato su tutti i bit 4 dei caratteri precedenti, e così via. Su ogni bit del blocco dati vengono eseguiti due controlli di parità: uno nella direzione orizzontale eseguito dal bit di parità del carattere ed un altro eseguito in direzione verticale dal carattere di controllo del blocco.

Spesso si indica col nome di bit di parità *orizzontale* il bit di parità del carattere, a cui viene anche attribuito il nome di bit di parità *trasversale*, di bit di parità *laterale* o di bit di parità della *riga*. Il bit di parità contenuto nel carattere di controllo del blocco viene chiamato spesso controllo di parità *longitudinale*, controllo di parità di *colonna* o controllo di parità *verticale*. Frequentemente, in letteratura si trova anche la locu-

zione *carattere di controllo di ridondanza longitudinale*, o più brevemente LRCC, dall'inglese *Longitudinal Redundancy Check Character*.

Il dispositivo di trasmissione aggiunge un carattere di controllo alla fine di ogni blocco di dati e trasmette i dati lungo la linea. Il destinatario calcola da parte sua quale dovrebbe essere il carattere di controllo del blocco sulla base dei caratteri via via ricevuti, e quindi confronta il carattere di controllo del blocco calcolato con il carattere di controllo del blocco ricevuto: se essi corrispondono, il destinatario dichiara valido il blocco.

Se così il blocco contenesse un solo bit errato, è possibile individuare l'errore perché si troverebbe errato sia il bit di controllo di parità del carattere (in cui il bit errato è contenu-

Fig. 6 - Individuazione degli errori con parità a due coordinate

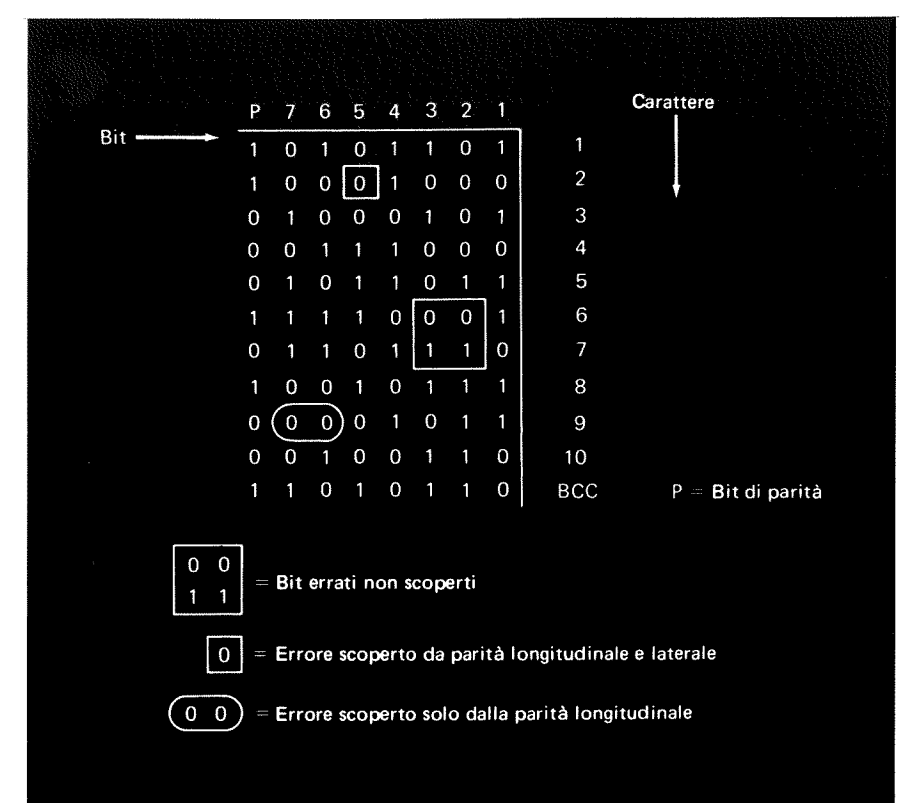
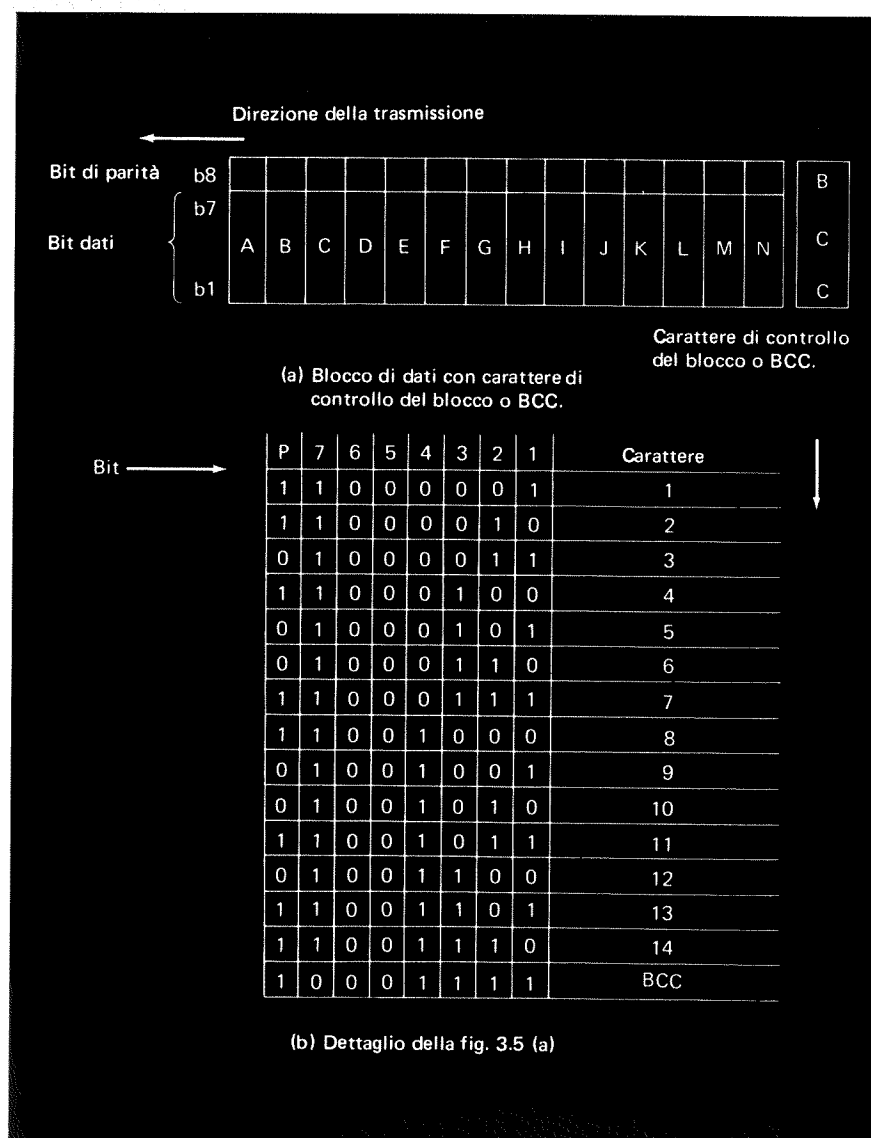


Fig. 5 - Controllo di parità a due coordinate



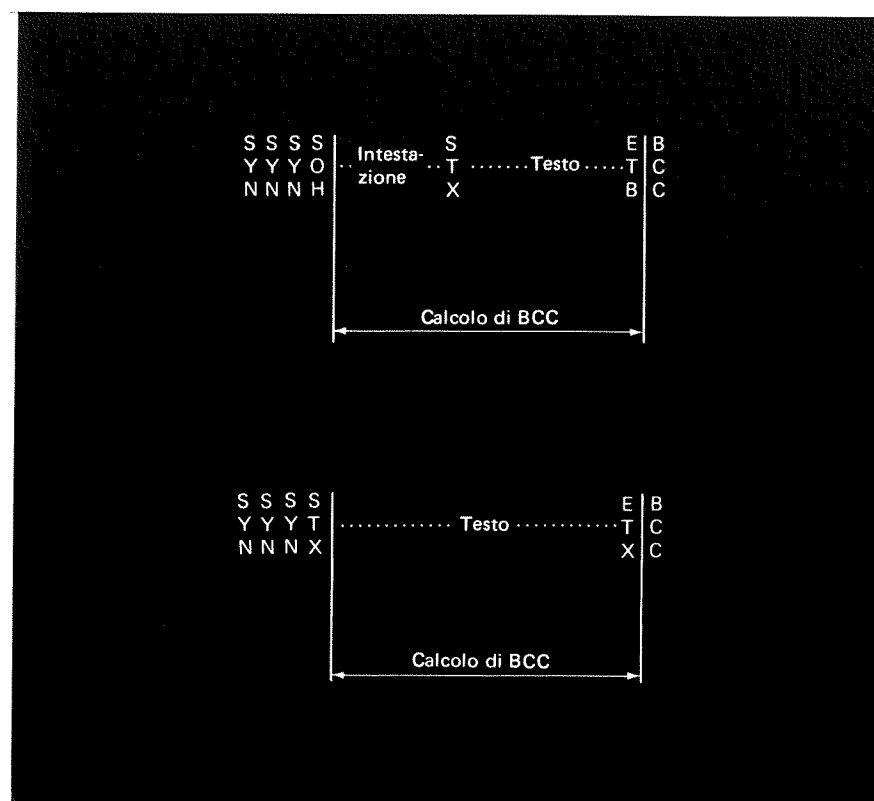
to) che il carattere di controllo del blocco. Ma se un carattere contenesse due errori, che tra di loro si compensano, si troverà che il bit di parità del carattere sarà corretto, ma non sarà corretto il carattere di controllo del blocco e si scoprirà che il blocco contiene errori anche se non sarà possibile individuare il carattere errato. Analogamente, se si verificassero due errori in due caratteri diversi e che si compensano trovandosi nella stessa posizione nei due caratteri, il carattere di controllo del blocco sarà corretto, ma risulterà errato il bit di parità dei caratteri errati, che si potranno quindi individuare. I calcoli delle due parità, orizzontale e verticale, sono controlli complementari e consentono di raggiungere una buona capacità d'individuazione degli errori del sistema. Tuttavia non sarà possibile scoprire due bit errati che si compensano vicendevolmente nell'ambito di un carattere e che si trovino ripetuti in un altro carattere dello stesso blocco di dati, ma la probabilità che questo capiti è molto bassa.

La fig. 6 presenta alcuni schemi d'errore che possono capitare. Il bit errato singolo contenuto nel carattere 2 verrebbe scoperto impiegando i controlli orizzontale e verticale: si

riuscirebbe ad individuare il doppio errore del carattere 9 con il controllo di parità verticale, mentre i quattro bit errati contenuti nei caratteri 6 e 7 non verrebbero scoperti dai due controlli. Di regola, se i bit errati giacciono sui vertici di un rettangolo, non vengono individuati con l'impiego congiunto dei controlli di parità e verticale ed orizzontale. Si possono analizzare matematicamente le condizioni nelle quali ci si trova ad operare e data la frequenza di errore che si può riscontrare su una data linea, determinare la lunghezza del blocco di trasmissione per ottimizzare il *throughput*.

È facile costruire il controllo di parità su due coordinate sia in software che in hardware, e questa è oggi la soluzione più frequentemente adottata. Il carattere di controllo del blocco viene costruito semplicemente eseguendo un OR esclusivo su tutti i caratteri precedenti. A seconda che si cominci a calcolare il carattere di controllo del blocco contando tutti gli 1 o tutti gli 0, avremo parità dispari o parità pari. La fig. 7 schematizza le regole per calcolare il carattere di controllo del blocco riportando due formati tipici di messaggio impiegati in un sistema di trasmissione sincrono. La sta-

zione trasmettente calcola il carattere di controllo del blocco iniziando dal primo carattere SOH - *Start Of Header* - o del carattere STX - *Start of Text* - trasmesso, carattere che però è compreso nel calcolo del controllo del blocco, quindi esegue un OR esclusivo su tutti i caratteri rimanenti fino a quando non trasmetterà per la prima volta il carattere ETB o il carattere ETX che saranno compresi nel calcolo: il carattere di controllo del blocco verrà trasmesso dopo lo ETB o lo ETX. In ricezione, il destinatario analizza i caratteri in arrivo per individuare il primo carattere di SOH o STX trasmesso. Appena riceve questo carattere, il destinatario comincia a calcolare da parte sua il carattere di controllo del blocco eseguendo un OR esclusivo su tutti i caratteri che seguono SOH o STX fino a quando verrà ricevuto il primo carattere ETX o ETB che verrà compreso nel calcolo. A questo punto il destinatario ha calcolato il «suo» carattere di controllo del blocco e lo confronta con il successivo carattere ricevuto dalla linea che è il carattere di controllo del blocco calcolato dal trasmettente. Se i due caratteri corrispondono, il blocco viene dichiarato valido. Se essi non corrispondono ci sarà segnalazione di



errore (il bit di parità del carattere sul carattere di controllo del blocco è, per inciso, il controllo di parità del carattere di controllo del blocco e non è il controllo di parità su tutti i bit di parità precedenti)).

Nel flusso di dati, dopo aver calcolato il carattere di controllo del blocco, si possono inserire dei caratteri SYN. Alcuni sistemi inseriscono dei caratteri SYN per guadagnar tempo, come riempitivo, poichè non sono sufficientemente veloci per mantenere la sincronizzazione tra i caratteri. Questi SYN vengono ignorati durante il calcolo del carattere di controllo del blocco: la maggior parte dei sistemi elimina i caratteri SYN dal flusso di dati e non li trasferisce al programma utente.

2. Controllo ciclico di ridondanza

Il controllo di parità a due coordinate è utile e consente di controllare in modo adeguato i sistemi di trasmissione dei caratteri, ma richiede un sensibile costo di elaborazione che è pari a un bit per ogni carattere, più un ulteriore carattere alla fine del blocco. Ora le case costruttrici tendono ad offrire apparecchi che consentono di trasmettere flussi di dati in forma binaria pura in cui non si deve necessariamente spezzare il

flusso di dati in caratteri singoli. Ciò impedisce d'applicare con facilità il controllo di parità a due coordinate.

Il controllo ciclico di ridondanza (o CRC, dall'inglese *Cyclic Redundancy Check*) ritrova un'applicazione sempre più diffusa a seguito dei progressi più recenti nella progettazione e costruzione dei circuiti hardware. E tale controllo è relativamente semplice da attuare impiegando i moderni circuiti integrati su larga scala.

Se invece di pensare che il flusso di dati in trasmissione è un flusso di singoli caratteri, lo considerassimo un lungo numero binario, potremmo dividerlo per un altro numero binario che chiameremo *costante*. Si tratta di una divisione in modulo due e non di una normale divisione aritmetica. Dividendo i *dati* per una *costante* si otterrà un *quoziente* e un *resto*: il resto verrà trasmesso lungo la linea immediatamente di seguito ai dati e, all'altro estremo, il destinatario eseguirà un'operazione analoga sui dati ricevuti, i quali saranno trattati come binario puro e saranno divisi per la stessa costante, calcolando quindi un resto e un quoziente. Il destinatario ignorerà il quoziente e confronterà il resto ottenuto con il resto ricevuto: se essi corri-

Fig. 7 - Calcolo del BCC o Block Control Character (carattere di controllo del blocco)

sponderanno, i dati verranno dichiarati validi; se invece saranno diversi, saranno considerati errati.

Come è indicato nella fig. 8, l'operazione viene eseguita in modo assai semplice, tutta in hardware. Il numero binario (cioè, in definitiva, i dati) viene trasmesso lungo la linea al destinatario e sarà letto contemporaneamente da un dispositivo hardware che eseguirà la divisione e calcolerà il resto che verrà poi trasmesso successivamente al destinatario che a sua volta avrà eseguito nel frattempo, all'altro estremo, le stesse operazioni nello stesso modo e confronterà il resto calcolato con il resto ricevuto.

6. Correzione dell'errore

Dopo aver scoperto un errore, si dovrà fare qualche cosa per correggerlo o quanto meno per minimizzarne l'effetto. Si possono gestire gli errori di trasmissione sostanzialmente in tre modi:

1. sostituzione del carattere errato,
2. invio della correzione dell'errore,
3. ritrasmissione.

1. Sostituzione con simbolo

In molti casi, quando in particolare i dati sono destinati ad essere letti da

un operatore, sarà sufficiente ricorrere al semplice controllo di parità. Se si scopre un errore di parità sul carattere, si sostituirà il carattere errato con qualche altro carattere. L'insieme dei caratteri (*set* di caratteri) Ascii contiene il carattere di controllo delle comunicazioni SUB che viene tradotto spesso in stampa o sullo schermo video con un punto interrogativo rovesciato: () o una sequenza con tre linee verticali (| | |). Se il messaggio contiene questo carattere esso verrà visualizzato e l'operatore sarà in grado d'individuare, e spesso correggere, l'errore.

2. Invio della correzione dell'errore

L'invio della correzione dell'errore richiede l'impiego di speciali codici

di trasmissione che contengono un'informazione ridondante sufficiente per individuare e correggere ogni errore all'estremo in ricezione. Per eseguire queste operazioni esiste un certo numero di codici che però non sono largamente utilizzati nelle applicazioni commerciali perchè comportano un elevato impiego di risorse del sistema, non giustificato dalla frequenza dei messaggi errati, in genere inferiore all'1%. Con tali frequenze è più conveniente che il destinatario che ha ricevuto un blocco con errori ne richieda la ritrasmissione.

3. Ritrasmissione

La ritrasmissione del blocco errato è di gran lunga il mezzo più usato per

la correzione degli errori. In sostanza, quando il destinatario scopre di aver ricevuto un blocco di dati contenente degli errori, si domanda alla stazione trasmittente di ripeterne la trasmissione, sperando che la seconda volta il blocco arrivi senza errori. La procedura di ritrasmissione dei blocchi impone un qualche dialogo tra la sorgente e il destinatario, che verrà controllato da un insieme di *procedure di controllo di linea*.

Nota: Il materiale di questo articolo è stato tratto dal libro dello stesso Autore «Sistemi di trasmissione dati e di teleelaborazione», collana dei Quaderni di Informatica, ed F. Angeli (1983).

L'elaboratore nello stoccaggio di gas naturale in pozzi esauriti

ERNESTO FIAMBERTI

Honeywell S.p.A.
Milano

1. Introduzione

Nell'impianto qui descritto in termini concettuali, realizzato recentemente nel nostro paese, si utilizza un giacimento esaurito di gas naturale come serbatoio polmone, cioè per immagazzinare temporaneamente del gas in eccesso rispetto al consumo. Il gas viene poi estratto e reimpresso nella rete di distribuzione nazionale per coprire le punte di carico richieste.

Tutto il processo di stoccaggio, erogazione e trattamento del gas viene gestito e monitorizzato mediante un complesso sistema di elaborazione basato su mini e microelaboratori.

2. L'impianto

Il giacimento esaurito usato come serbatoio di stoccaggio è formato da 2 gruppi di pozzi, che fanno capo ognuno ad un collettore di invio del gas diretto ai pozzi (stoccaggio) o di prelievo del gas di reimmissione nella rete (distribuzione).

Per ottenere la portata di gas desiderata, ogni pozzo è dotato di una valvola ad apertura variabile.

Prima della reimmissione del gas nella rete di distribuzione, ne viene fatto un trattamento di disidratazione mediante glicol.

È anche presente un impianto di rigenerazione del glicol di disidratazione.

Il trattamento del gas prima della sua reimmissione nella rete di distribuzione, avviene mediante dodici colonne di disidratazione e sei im-

pianti di rigenerazione del glicol di disidratazione.

Le due fasi di funzionamento dell'impianto sono: stoccaggio ed erogazione.

In fase di stoccaggio si chiude tutto l'impianto di trattamento ed il gas dalle linee di stoccaggio viene convogliato ai vari pozzi.

In fase di erogazione, il gas in uscita dai pozzi subisce una prima separazione per gravità dai liquidi che fuoriescono dal pozzo.

Prima della immissione nel collettore, al gas viene unito del glicol (glicol di iniezione) onde evitare che nei punti di salto di pressione si formino degli idrati. Successivamente avviene il trattamento nelle colonne di disidratazione, dopo aver subito un riscaldamento in un forno.

Un regolatore di portata modula una valvola per ottenere la portata di gas desiderata all'ingresso di ogni colonna.

Sul fondo colonna viene raccolto del glicol che, dopo un processo di rigenerazione, viene utilizzato come glicol di iniezione.

Il processo di disidratazione nella colonna avviene in quanto il gas umido in corrente ascensionale si scontra nella testa colonna con altro glicol, che ne elimina l'umidità grazie alla proprietà igroscopica.

Anche questo glicol di disidratazione subisce un processo successivo di rigenerazione.

Infine il gas disidratato, attraverso due linee di erogazione, viene im-

messo nella rete di distribuzione.

Le apparecchiature di controllo dell'impianto sono costituite da due sottosistemi:

- sottosistema di acquisizione e restituzione dati (TDC 2000 Honeywell), dotato di una "operator station" che consente di controllare il processo, sia pure con una degradazione delle prestazioni, in assenza di supervisore vero e proprio;
- sottosistema duale di elaborazione e controllo costituito da un microelaboratore di front-end che, oltre a colloquiare col sottosistema di acquisizione e restituzione dati, esegue delle preelaborazioni sui dati acquisiti, e da due minielaboratori che eseguono le funzioni di elaborazione e controllo, e interagiscono con gli operatori.

3. Funzioni realizzate

Attraverso le funzionalità combinate della rete di acquisizione dati e dei programmi effettuati dall'unità centrale, gli operatori hanno accesso sui terminali video alle seguenti informazioni e comandi:

- la visualizzazione dell'elenco (ordinato cronologicamente) degli allarmi e degli eventi di cui non hanno ancora preso visione;
- la visualizzazione dell'elenco degli allarmi di cui hanno già preso visione e che sono ancora in atto;
- la visualizzazione di schemi a struttura predefinita (per mezzo di un programma configuratore) nei quali vengono presentati e

continuamente aggiornati valori e/o stato di variabili di processo (valori di misure, stati di dispositivi, stati di allarme, ore di funzionamento);

- la hard copy su stampante degli schemi visualizzati;
- la stampa di sommari con i dati di stoccaggio e/o erogazione del gas (correnti o relativi a periodi precedenti);
- l'invio di comandi elementari ai dispositivi per mezzo di una procedura di selezione sia del dispositivo sia del comando guidato dal sistema e caratterizzato dallo spostamento mediante tasti funzionali di un cursore in posizioni predefinite (solo nella modalità "supervisione operativa").

Al fine di rendere disponibili agli operatori le informazioni sopra elencate, sui segnali analogici e digitali acquisiti dal TDC 2000 vengono eseguiti dei controlli di attendibilità e viene ricavato lo stato di allarme per confronto con valori limite di ciascuna misura.

Vengono calcolati dei valori e degli stati non direttamente acquisibili e, confrontando stati attuali di misure con quelli acquisiti o calcolati alla precedente elaborazione, vengono generati opportuni messaggi e segnalazioni presentati automaticamente su stampante e, a richiesta, su terminale video.

Vengono eseguiti anche programmi di ausilio manutentivo per gli impianti come ad esempio il conteggio delle ore di funzionamento dei pozzi, delle pompe e dei forni.

Per la gestione da parte dell'operatore, l'impianto è stato suddiviso nelle seguenti parti:

- 2 gruppi (cluster) con una quindicina di pozzi ciascuno (in futuro vi sarà una espansione, già prevista nel sistema, a 4 cluster con un totale di circa 50 pozzi);
- 1 collettore (manifold) di ingresso;
- 2 linee di erogazione;
- 12 colonne di disidratazione;
- 1 impianto di rigenerazione e di invio ai cluster di glicol di iniezione;
- 3 impianti di rigenerazione del glicol di disidratazione;
- servizi di centrale.

Ciascuna di queste parti di impianto può essere controllata dal sistema (compatibilmente con le modalità operative del resto dell'impianto) in tre modi possibili:

- supervisione non operativa
- supervisione operativa
- conduzione automatica.

Il passaggio ai diversi modi operativi viene attuato su richiesta degli operatori. L'accettazione del passaggio dipende però dalle modalità operative delle altre parti di impianto.

La supervisione non operativa permette al sistema di conoscere lo stato di quella parte di impianto senza poter effettuare peraltro alcuna modifica.

La supervisione operativa permette di attuare solo degli interventi elementari sull'impianto, tramite richiesta dell'operatore.

La conduzione automatica è invece basata sulla esecuzione periodica di sequenze, che possono anche essere attivate da eventi.

Dette sequenze possono essere associate ad una parte di impianto o a tutto l'impianto per quelle parti dichiarate in conduzione automatica.

Ogni parte di impianto ha determinati dispositivi che possono essere attivati dalla sequenza. I dispositivi possono essere anche considerati vincolati, durante la conduzione automatica, allo stato di inattività per motivi quali condizioni di allarme, richieste esterne al programma, fallimento di più comandi consecutivi. Lo stato di vincolo può essere rimosso solo dagli operatori tramite una funzione interattiva.

Tra le sequenze sono di particolare importanza la sequenza di allarmi e dell'invio e controllo comandi.

La sequenza allarmi viene eseguita ad ogni variazione di condizioni di allarme associate a parti di impianto in conduzione automatica dopo la rimozione di cause di vincolo e ad ogni passaggio di parti di impianto in conduzione automatica.

Per ogni condizione di allarme in atto non ancora trattata dopo la sua comparsa o dopo la rimozione di cause di vincolo, vengono imposti gli stati di vincolo ai dispositivi associati e viene richiesta alla sequenza

di invio e controllo comandi l'eventuale esecuzione di comandi di emergenza volti a mettere gli impianti in condizioni di sicurezza. Quando le condizioni di allarme richiedono un'analisi o un'azione più complessa, la sequenza allarmi può rimandare ad altre sequenze queste operazioni limitandosi a chiederne l'esecuzione.

La sequenza di invio e controllo comandi riceve da tutte le altre sequenze le richieste di esecuzione di comandi e le attua, verificando l'effetto delle manovre e ripetendole fino a tre volte in caso di fallimento.

La sequenza può ricevere richieste di invio di comandi operativi o di emergenza.

La distinzione non riguarda i comandi, ma il motivo per cui vengono richiesti (ad es. un comando di chiusura di un pozzo è operativo se dovuto a una richiesta di riduzione di portata, di emergenza se dovuto a una condizione di allarme).

I comandi di emergenza vengono sempre eseguiti, quelli operativi soltanto se il dispositivo associato non è vincolato e il suo stato è diverso da quello che il comando gli farebbe assumere.

Gli obiettivi da raggiungere durante la conduzione automatica del processo di erogazione (il più complesso) sono:

- erogazione verso la rete nazionale della portata desiderata;
- utilizzo del numero minimo necessario di colonne di disidratazione;
- funzionamento delle colonne il più possibile vicino al punto ottimale (portata di uscita pari a 80% della portata massima);
- utilizzazione del maggior numero di pozzi alla minima apertura consentita dalla portata desiderata;
- mantenimento del valore ottimale della pressione del gas nel manifold di ingresso.

A questo scopo si fissano in funzione delle richieste degli operatori due set-point, uno relativo alla portata in uscita verso la rete nazionale, l'altro coincidente con il valore ottimale di pressione nel manifold.

I due set-point vengono inseguiti in modo del tutto indipendente.

Per raggiungere quello di portata nel rispetto dei criteri sopra elencati, si assegna a ciascuna linea di erogazione un set-point di portata (sequenza di distribuzione della portata alle linee di erogazione), si distribuisce la portata assegnata a ciascuna linea tra le sue colonne (sequenze di distribuzione della portata alle colonne) e si calcola il set-point da fornire al loop di regolazione di ciascuna colonna realizzato nel Basic Controller del TDC 2000 (sequenze di controllo della portata della colonna).

Per raggiungere invece il set-point di pressione si ricava da dati ottenuti superiormente la variazione di aper-

tura complessiva di pozzi (che sono ad apertura modulabile) in funzione del valore ottimale di pressione e dello scostamento tra il valore attuale e quello ottimale; la realizzazione di questa variazione viene distribuita ai cluster (sequenza di distribuzione del carico ai cluster).

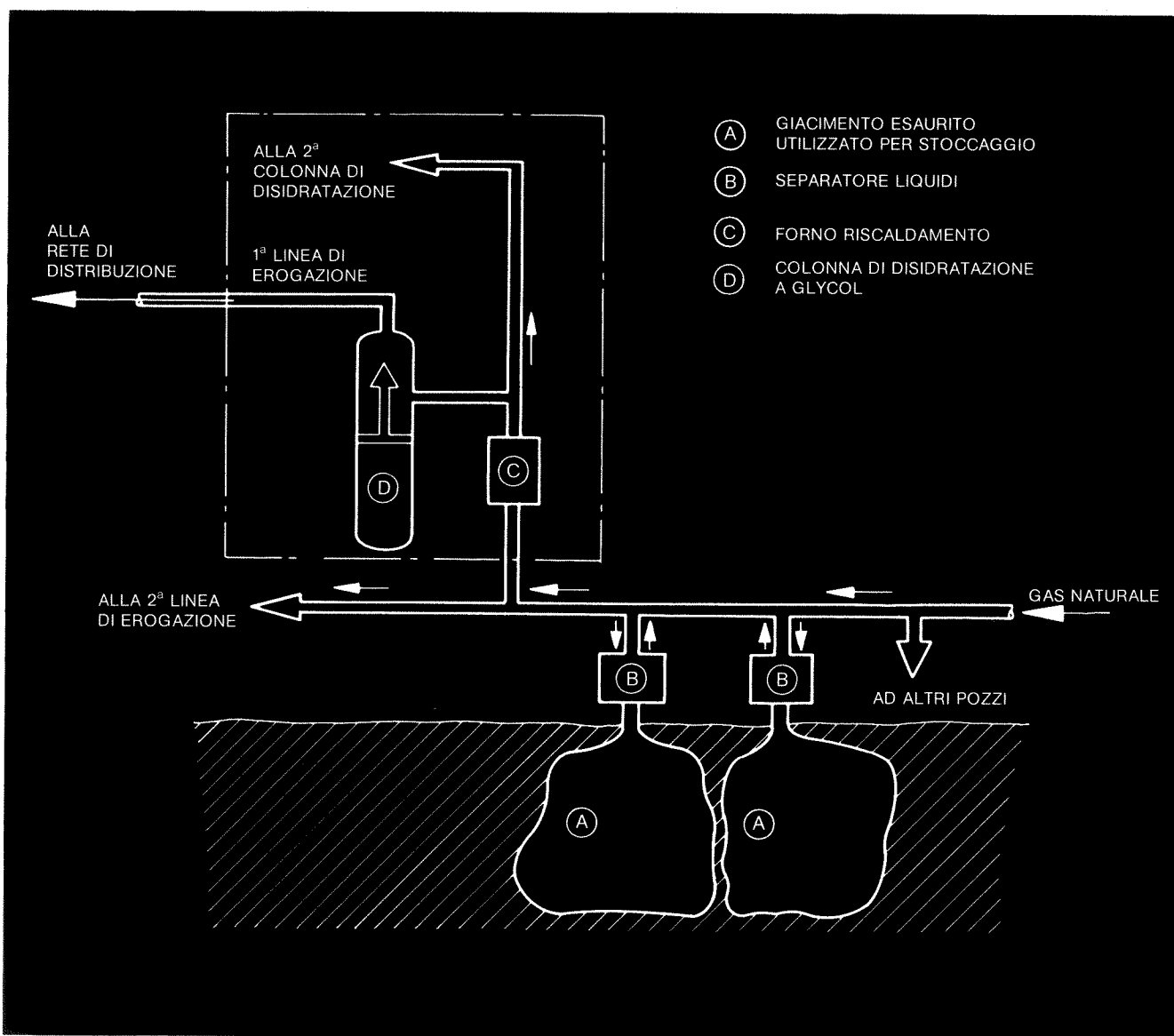
Successivamente viene calcolato e raggiunto il livello di apertura di ciascun pozzo (sequenze di distribuzione di carico ai pozzi).

Nelle normali condizioni di funzionamento il set-point di pressione sarà mantenuto ad un valore fisso (quello ottimale), mentre il set-point di portata verrà variato in fun-

zione delle esigenze della rete di distribuzione.

Quando viene richiesta una variazione di portata le sequenze di controllo dell'erogazione in centrale provvedono a soddisfarla senza curarsi delle ripercussioni a monte della centrale (accumulo o rarefazione del gas). Ciò produrrà una variazione della pressione nel manifold e l'intervento delle sequenze di controllo relative che tenderanno a riportarla al valore ottimale modulando opportunamente l'apertura dei pozzi.

Schema semplificato dell'impianto di stoccaggio, erogazione e trattamento di gas naturale.



Macchine calcolatrici e intelligenza^(*)

A.M. Turing

1. Il gioco dell'imitazione

Mi propongo di considerare la questione: "Possono pensare le macchine?" Si dovrebbe cominciare col definire il significato dei termini "macchina" e "pensare". Le definizioni potrebbero essere elaborate in modo da riflettere il più possibile l'uso normale delle parole, ma questo atteggiamento è pericoloso. Se il significato delle parole "macchina" e "pensare" deve essere trovato esaminando le parole stesse attraverso il loro uso comune è difficile sfuggire alla conclusione che tale significato e la risposta alla domanda "possono pensare le macchine?" vadano ricercati in una indagine statistica del tipo delle inchieste Gallup. Ciò è assurdo. Invece di tentare una definizione di questo tipo sostituirò la domanda con un'altra, che le è strettamente analoga e che è espressa in termini non troppo ambigui.

La nuova forma del problema può essere descritta nei termini di un gioco, che chiameremo "il gioco dell'imitazione". Questo viene giocato da tre persone, un uomo (A), una donna (B) e l'interrogante (C), che può essere dell'uno o dell'altro sesso. L'interrogante viene chiuso in una stanza, separato dagli altri

due. Scopo del gioco per l'interrogante è quello di determinare quale delle altre due persone sia l'uomo e quale la donna. Egli le conosce con le etichette X e Y, e alla fine del gioco darà la soluzione "X è A e Y è B" o la soluzione "X è B e Y è A". L'interrogante può far domande di questo tipo ad A e B:

C: "Vuol dirmi X, per favore, la lunghezza dei propri capelli?"

Ora supponiamo che X sia in effetti A, quindi A deve rispondere. Scopo di A nel gioco è quello di ingannare C e far sì che fornisca una identificazione errata. La sua risposta potrebbe perciò essere: "I miei capelli sono tagliati a la garçonne, ed i più lunghi sono di circa venticinque centimetri". Le risposte, in modo tale che il tono di voce non possa aiutare l'interrogante, dovrebbero essere scritte, o, meglio ancora, battute a macchina. La soluzione migliore sarebbe quella di avere una telescrivente che mettesse in comunicazione le due stanze. Oppure le domande e risposte potrebbero essere ripetute da un intermediario. Scopo del gioco, per il terzo giocatore (B), è quello di aiutare l'interrogante. La migliore strategia per lei è probabilmente quella di dare risposte veritiere. Essa può anche aggiungere alle sue risposte frasi come "sono io la donna, non dargli ascolto!", ma ciò approderà a nulla dato che anche l'uomo può fare affermazioni analoghe. Poniamo ora la domanda: "Che cosa accadrà se una macchina prenderà il posto di A nel gioco?" L'in-

terrogante darà una risposta errata altrettanto spesso di quando il gioco viene giocato tra un uomo e una donna? Queste domande sostituiscono quella originale: "Possono pensare le macchine?"

2. Critica del nuovo problema

Come si potrebbe domandare "Quale è la risposta alla domanda nella sua nuova formulazione?", così si potrebbe anche chiedere "La nuova domanda merita di ricevere una risposta?" Esamineremo subito l'ultimo quesito, tagliando corto in tal modo ad un regresso all'infinito.

Il nuovo problema ha il vantaggio di tirare una linea di separazione abbastanza netta tra le capacità fisiche e quelle intellettuali di un uomo. Nessun ingegnere o chimico pretende di essere capace di produrre un materiale che non si possa distinguere dalla pelle umana. Può darsi che un giorno questo possa essere fatto, ma perfino supponendo disponibile una invenzione siffatta riterremmo che val poco la pena cercare di rendere più umana una "macchina pensante" rivestendola a questo modo di carne artificiale. La forma nella quale abbiamo posto il problema riflette questo fatto, nella condizione che impedisce all'interrogante di vedere o toccare i due competitori, e di udire le loro voci. Altri vantaggi del criterio proposto possono essere messi in luce da domande e risposte campione. Per esempio:

Domanda: Mi scriva, per favore, un sonetto sul tema *Forth Bridge*.

(*) Questo scritto viene qui proposto in chiave storica, essendo apparso originariamente sulla rivista "Mind" nel 1950. Il suo autore, Alan Mathison Turing (1912-1954), una delle menti più brillanti nel campo della logica matematica, ha fornito contributi fondamentali alla teoria delle macchine calcolatrici.

Risposta: Non faccia affidamento su di me per questo. Non ho mai saputo scrivere poesie.

Domanda: Sommi 34957 a 70764.

Risposta (pausa di circa trenta secondi e poi la risposta): 105721.

Domanda: Giuoca a scacchi?

Risposta: Sì.

Domanda: Ho il Re in *ei* e nessun altro pezzo. Lei ha solo il Re in *c3* ed una Torre in *h8*. Tocca a lei. Che mossa giuoca?

Risposta (dopo una pausa di quindi secondi): Torre in *h1*, matto.

Il metodo delle domande e risposte sembra essere adatto per introdurre nell'esame quasi ogni campo della conoscenza umana che desideriamo. Non desideriamo penalizzare la macchina per la sua incapacità di brillare in un concorso di bellezza, né penalizzare un uomo perché perde una corsa contro un aeroplano. Le condizioni del nostro giuoco rendono irrilevanti queste incapacità. I "testimoni" possono vantarsi quanto vogliono, se lo considerano opportuno, della loro bellezza, forza ed eroismo, ma l'interrogante non può chiedere dimostrazioni pratiche.

Il giuoco può forse essere criticato sulla base del fatto che le possibilità sono troppo nettamente a sfavore della macchina. Se l'uomo dovesse cercare di fingere di essere la macchina farebbe certamente una figura molto brutta. Sarebbe tradito immediatamente dalla sua lentezza e imprecisione nell'aritmetica. Non possono forse le macchine comportarsi in qualche maniera che dovrebbe essere descritta come pensiero ma che è molto differente da quanto fa un uomo? Questa obiezione è molto forte, ma come minimo possiamo dire che se, ciononostante, una macchina può essere costruita in modo da giocare il giuoco dell'imitazione soddisfacentemente, non abbiamo bisogno di tenerne conto.

Si potrebbe insistere che giocando il giuoco dell'imitazione la migliore strategia per la macchina potrebbe non essere forse l'imitazione del comportamento di un uomo. Può anche darsi, ma non credo che si possa dare grande peso ad una possibilità del genere. In ogni caso non

è nostra intenzione qui esaminare la teoria del giuoco, e sarà dato per scontato che la migliore strategia per la macchina sia quella di provare a formulare le risposte che sarebbero date istintivamente da un uomo.

3. Le macchine interessate al giuoco

La domanda che abbiamo posto nel paragrafo 1 non sarà del tutto definita fino a quando non avremo specificato che cosa intendiamo con la parola "macchina". Naturalmente sarebbe nostro desiderio che ogni tipo di tecnica ingegneristica potesse essere usata nella costruzione delle nostre macchine. Sarebbe pure nostro desiderio permettere che un ingegnere o una squadra di ingegneri potesse costruire una macchina che funzionasse, ma i cui metodi di operare non potessero essere descritti in maniera soddisfacente dai suoi costruttori in quanto essi hanno applicato metodi largamente sperimentali. Infine, vorremmo escludere dal concetto di macchina gli uomini nati nel modo normale. È difficile adattare le definizioni in modo tale da soddisfare queste tre condizioni. Si potrebbe per esempio richiedere che la squadra di ingegneri sia composta tutta da ingegneri dello stesso sesso, ma questo non sarebbe del tutto soddisfacente, perché è probabilmente possibile dar vita ad un individuo completo da una singola cellula della pelle, poniamo, di un uomo. Sul piano della tecnica biologica un risultato del genere sarebbe tale da meritare la lode più alta, ma non saremmo inclini a considerarlo un caso di "costruzione di una macchina pensante". Questo ci obbliga ad abbandonare l'esigenza di permettere l'uso di qualsiasi tipo di tecnica. Siamo tanto più pronti a far questo in vista del fatto che l'attuale interesse alle "macchine pensanti" è stato destato da un particolare tipo di macchina, chiamato correntemente "calcolatore elettronico" o "calcolatore numerico". Seguendo questo indirizzo permetteremo soltanto ai calcolatori numerici di prendere parte al nostro giuoco.

Questa restrizione a prima vista appare molto drastica. Cercherò di dimostrare che in realtà non è così. Si

rende perciò necessario dare, in breve, notizia della natura e delle proprietà di questi calcolatori.

Si potrebbe anche sostenere che l'identificazione delle macchine con i calcolatori numerici, come pure il nostro criterio per il concetto di "pensare", sarebbe insoddisfacente se (contrariamente a quanto credo) risultasse che i calcolatori numerici sono incapaci di fare una buona figura nel giuoco. C'è già un buon numero di calcolatori numerici in funzione e si potrebbe chiedere: "Perché non tentare immediatamente l'esperimento? Sarebbe facile soddisfare le condizioni del giuoco. Potrebbe essere usato un certo numero di esaminatori, e potrebbero essere compilate statistiche per mostrare in quale proporzione è stata fornita l'identificazione giusta". La risposta è che non ci stiamo chiedendo se tutti i calcolatori numerici potrebbero far buona figura nel giuoco né se potrebbero far buona figura nel giuoco i calcolatori attualmente disponibili, ma se siano immaginabili calcolatori che potrebbero farla. Ma questa è solo una risposta per tagliar corto alle discussioni. Considereremo più tardi questa questione sotto una luce diversa.

4. Calcolatori numerici

L'idea che sta alla base dei calcolatori numerici può essere spiegata dicendo che queste macchine sono costruite per compiere qualsiasi operazione che possa essere compiuta da un calcolatore umano. Si suppone che il calcolatore umano segua regole fisse; egli non ha l'autorità di deviare da esse in alcun dettaglio. Possiamo supporre che queste regole siano fornite da un libro, che viene modificato ogni volta che egli viene adibito ad un nuovo lavoro. Egli ha pure una riserva illimitata di carta sulla quale fare i suoi calcoli. Può anche compiere le sue moltiplicazioni ed addizioni con una calcolatrice da tavolo, ma questo non è importante.

Se usiamo la spiegazione data sopra come una definizione rischieremo di cadere in un circolo vizioso. Evitiamo questo rischio dando una in-

dicazione dei metodi attraverso i quali l'effetto desiderato viene raggiunto. Un calcolatore numerico può essere normalmente considerato composto di tre parti:

- 1) memoria,
- 2) complesso operativo,
- 3) governo.

La memoria è un deposito di informazioni, e corrisponde alla carta del calcolatore umano, sia che si tratti della carta sulla quale egli fa i suoi calcoli, sia di quella sulla quale è stampato il suo libro di regole. Per quella parte dei calcoli che il calcolatore umano compie con il suo cervello, una parte di questo deposito corrisponderà alla sua memoria.

Il complesso operativo è la parte che compie le varie operazioni singole che un calcolo comporta. Quali saranno queste singole operazioni dipenderà dalle diverse macchine. Di solito possono essere compiuti calcoli piuttosto lunghi come "moltiplicare 3.540.675.445 per 7.076.345.687", ma in alcune macchine sono possibili soltanto operazioni molto semplici del tipo "scrivere zero".

Abbiamo fatto presente che il "libro delle regole" fornito al calcolatore umano è sostituito nella macchina da una parte della memoria. Si chiama allora "tavola delle istruzioni". È compito del "governo" controllare che queste istruzioni siano eseguite correttamente e nell'ordine giusto. Il governo è costruito in maniera tale che questo avviene necessariamente.

Le informazioni contenute nella memoria sono comunemente suddivise in sezioni di dimensioni abbastanza ridotte. In una macchina, per esempio, una sezione potrebbe consistere di dieci cifre decimali. Vengono assegnati dei numeri alle parti della memoria nelle quali le varie sezioni di informazione vengono immagazzinate, in una qualche maniera sistematica. Una istruzione tipica potrebbe dire:

"Sommare il numero immagazzinato nella posizione 6809 a quello contenuto nella posizione 4302 e riportare il risultato in quest'ultima cella di memoria".

Inutile dire che non sarebbe necessario esprimersi in inglese nei confronti della macchina. L'ordine sarebbe più probabilmente codificato in una forma del tipo 6.809.430.217. Qui 17 indica quale delle possibili operazioni deve essere compiuta sui due numeri. In questo caso l'operazione è quella descritta sopra, cioè: "Sommare il numero..." Si sarà notato che l'istruzione comprende 10 cifre e forma così un blocchetto di informazione, cosa molto conveniente. Il governo normalmente farà in modo che le istruzioni siano eseguite secondo l'ordine delle posizioni nelle quali sono immagazzinate, ma occasionalmente si possono incontrare istruzioni del tipo:

"obbedire alla istruzione immagazzinata nella posizione 5606, e continuare"

oppure

"se la posizione 4505 contiene uno zero obbedire poi alla istruzione 6707, altrimenti continuare."

Istruzioni di questi ultimi tipi sono molto importanti perché rendono possibile ripetere più volte una sequenza di operazioni fin quando siano soddisfatte determinate condizioni, ma obbedendo, in questo modo, non a nuove istruzioni per ogni ripetizione dell'operazione, ma sempre alle stesse. Per fare una analogia di carattere familiare, supponiamo che la mamma voglia che Tommy passi dal calzolaio ogni mattina mentre va a scuola per vedere se le sue scarpe sono pronte. Essa può chiederglielo di nuovo tutte le mattine. Oppure può appendere una volta per tutte un avviso nell'ingresso, che egli vedrà ogni volta che esce per andare a scuola, e che gli ricorda di occuparsi delle scarpe; e può anche distruggere l'avviso quando lui le ha ritirate.

Il lettore deve accettare come un dato che i calcolatori numerici possono essere costruiti, ed in effetti sono stati costruiti, secondo i principi che abbiamo descritto, e che essi possono in effetti imitare molto da vicino le azioni di un calcolatore umano.

Il libro delle regole di cui abbiamo attribuito l'uso al nostro calcolatore umano è naturalmente una comoda definizione. Il calcolatore umano, in

realtà, si ricorda di quello che ha da fare. Se si vuole che una macchina imiti il comportamento di un calcolatore umano in una serie di operazioni complesse si deve chiedere all'uomo come vanno compiute quelle operazioni e quindi tradurre la risposta nella forma di una tavola di istruzioni. Il costruire tavole di istruzioni è comunemente descritto come "programmare". "Programmare una macchina in modo che compia l'operazione A" significa inserire nella macchina una tavola di istruzioni appropriate, in modo tale che essa compia l'operazione A.

Una variante interessante dell'idea di calcolatore numerico è quella di "calcolatore numerico con un elemento casuale". Queste macchine contengono istruzioni che comportano operazioni come gettare un dado o qualche processo elettronico equivalente; una istruzione del genere potrebbe essere, per esempio, "gettare il dado e memorizzare il numero che risulta nella posizione 1000". Alcune volte macchine del genere sono descritte come fornite di un libero arbitrio (anche se io non userei una frase del genere). Normalmente risulta impossibile determinare semplicemente osservando una macchina se essa contiene un elemento casuale, perché un effetto analogo può essere prodotto con un espediente, facendo dipendere le scelte, ad esempio, dalle successive cifre decimali di π .

La maggior parte dei calcolatori numerici esistenti hanno soltanto una memoria finita. Non c'è difficoltà teorica nell'idea di un calcolatore con una memoria illimitata. Naturalmente a ciascun istante potrà esserne usata soltanto una parte finita. Allo stesso modo può esserne stata costruita soltanto una parte finita, ma possiamo immaginarci successive aggiunte a seconda delle esigenze. Calcolatori del genere hanno un interesse teorico particolare e li chiameremo calcolatori a capacità infinita.

L'idea di un calcolatore numerico è piuttosto vecchia. Charles Babbage, professore di matematica a Cambridge dal 1828 al 1839, progettò una macchina del genere, da lui battezzata *macchina analitica*, che però

non fu mai completata. Sebbene Babbage avesse tutte le idee essenziali, la sua macchina a quel tempo non lasciava prevedere prospettive molto attraenti. Le velocità disponibili sarebbero state certamente superiori a quelle di un calcolatore umano, ma all'incirca 100 volte inferiori a quelle della macchina di Manchester, una delle macchine moderne più lente. La memoria sarebbe stata puramente meccanica, a base di ingranaggi e schede. Il fatto che la macchina analitica di Babbage dovesse essere nelle intenzioni del suo ideatore interamente meccanica ci aiuterà a liberarci di una superstizione. Si è data spesso importanza al fatto che i moderni calcolatori numerici sono elettrici, e che è elettrico anche il sistema nervoso. Dato che la macchina di Babbage non era elettrica, e dato che tutti i calcolatori numerici sono in un certo senso equivalenti, è evidente che questo uso della elettricità non può avere importanza teorica. Naturalmente l'elettricità interviene ogni volta che occorrono segnalazioni rapide, quindi non c'è da sorprendersi se la troviamo in entrambi questi casi. Nel sistema nervoso i fenomeni chimici sono almeno altrettanto importanti di quelli elettrici. In alcuni calcolatori il sistema di memorizzazione è prevalentemente acustico. La caratteristica di servirsi della elettricità dà luogo quindi soltanto ad analogie molto superficiali. Desiderando trovare somiglianze tra calcolatori e sistema nervoso dovremmo piuttosto interessarci delle analogie matematiche di funzionamento.

5. Universalità dei calcolatori numerici

I calcolatori numerici esaminati nell'ultimo paragrafo possono essere classificati tra le "macchine a stati discreti", cioè tra quelle che si muovono a salti o scatti improvvisi da uno stato ben definito ad un altro. Questi stati sono abbastanza differenti perchè si possa ignorare la possibilità di confusione tra di essi. Strettamente parlando non esistono macchine del genere. In realtà ogni cosa si muove con continuità. Ma ci sono molti tipi di macchine che pos-

sono vantaggiosamente essere viste come macchine a stati discreti.

Per esempio, considerando gli interruttori di un sistema di illuminazione è comodo supporre che ogni interruttore sia decisamente chiuso o aperto. Necessariamente esistono posizioni intermedie, ma per la maggior parte degli scopi possono essere trascurate. Come esempio di macchina a stati discreti possiamo considerare una ruota che scatti girando di 120 gradi una volta al secondo, ma che può essere fermata da una leva azionata dall'esterno; inoltre vi sia una lampada che si accende in una delle posizioni della ruota. Una tale macchina può essere descritta in forma astratta come segue. Lo stato interno della macchina (che è indicato dalla posizione della ruota) può essere q_1 , q_2 , o q_3 . C'è un segnale di ingresso i_0 o i_1 (posizione della leva). Ad ogni istante lo stato interno è determinato dall'ultimo stato e dal segnale di ingresso secondo la tabella:

		Ultimo stato		
		q_1	q_2	q_3
Ingresso	i_0	q_2	q_3	q_1
	i_1	q_1	q_2	q_3

I segnali in uscita, sole indicazioni esternamente visibili dello stato interno (la luce), sono descritti dalla tabella:

Stato	q_1	q_2	q_3
Uscita	u_0	u_0	u_1

Questo è un esempio rappresentativo delle macchine a stati discreti. Esse possono essere descritte da tavole del genere purché abbiano soltanto un numero finito di stati possibili.

Sembrerà che, dati lo stato iniziale della macchina ed i segnali di ingresso, sia sempre possibile predire tutti gli stati successivi. Questo ricorda l'ipotesi di Laplace che dallo stato completo dell'universo ad un momento dato, descritto mediante la posizione e le velocità di ogni particella, sia possibile predire tutti gli stati futuri. La predisposizione che stiamo esaminando è, tuttavia, mol-

to più vicina alla realizzazione pratica di quella formulata da Laplace. Il sistema dell'"universo come un tutto" è tale che errori molto piccoli nelle condizioni iniziali possono avere effetti disastrosi in un momento successivo. Lo spostamento di un singolo elettrone per un miliardesimo di centimetro, ad un momento dato potrebbe significare la differenza tra due avvenimenti molto diversi, come l'uccisione di un uomo un anno dopo a causa di una valanga o la sua salvezza. È una proprietà essenziale dei sistemi meccanici che abbiamo chiamati "macchine a stati discreti", che questo fenomeno non si verifica. Perfino quando consideriamo le attuali macchine concrete, in luogo di quelle idealizzate, una conoscenza sufficientemente precisa del loro stato in un dato momento porta ad una conoscenza sufficientemente precisa del loro stato ad un qualsiasi numero di passi successivi. Come abbiamo fatto presente, i calcolatori numerici rientrano nella classe delle macchine a stati discreti. Ma il numero di stati dei quali una macchina del genere è capace è di solito enormemente elevato. Per esempio per la macchina che lavora in questo momento a Manchester è di circa 2^{165000} , cioè di circa 10^{50000} . Si può fare un paragone tra questa macchina e il nostro esempio della ruota a scatti descritta sopra, che aveva tre stati. Non è difficile vedere perchè il numero degli stati debba essere così elevato. Il calcolatore implica una memoria corrispondente alla carta usata dal calcolatore umano. Deve essere possibile inserire nelle memorie qualsiasi combinazione di simboli che potrebbe essere stata scritta sulla carta. Per semplicità supponiamo che vengano usati come simboli soltanto le cifre da 0 a 9. Le variazioni di grafia vengono trascurate. Supponiamo che il calcolatore abbia a disposizione 100 fogli di carta contenenti ognuno 50 righe con spazio sufficiente per 30 cifre. Il numero degli stati allora diviene $10^{100 \times 50 \times 30}$, cioè 10^{150000} . Questo è all'incirca il numero degli stati di tre macchine di Manchester messe insieme. Il logaritmo in base due del numero di stati è chiamato di solito la "capacità di memorizzazione" della mac-

china. Così la macchina di Manchester possiede una capacità di memorizzazione di circa 165000 e la macchina a ruota del nostro esempio di circa 1,6. Se due macchine sono messe insieme le loro capacità devono essere addizionate per ottenere la capacità della macchina che ne risulta.

Data la tavola corrispondente allo stato di una macchina a stati discreti, è possibile predire cosa essa farà. Non c'è alcuna ragione per cui questo calcolo non possa essere compiuto facendo uso di un calcolatore numerico. Purché esso possa venir eseguito con una velocità sufficiente, il calcolatore numerico potrebbe imitare il comportamento di qualsiasi macchina a stati discreti. Il giuoco dell'imitazione potrebbe allora essere giocato tra la macchina in questione (nella parte di B) e il calcolatore numerico (nella parte di A) e l'interrogante non sarebbe capace di distinguerli. Naturalmente il calcolatore numerico dovrebbe avere una memoria adeguata e funzionante a velocità abbastanza alta. Inoltre bisognerebbe programmarlo di nuovo per ogni nuova macchina che si desidera fargli imitare.

Questa speciale proprietà dei calcolatori numerici, cioè che essi possono imitare ogni macchina a stati discreti, si può descrivere dicendo che essi sono macchine *universali*. L'esistenza di macchine con questa proprietà ha la conseguenza importante che, a parte considerazioni di velocità, non è necessario progettare varie macchine differenti per compiere processi differenti di calcolo. Questi possono essere compiuti tutti con un solo calcolatore numerico programmato nella forma adatta caso per caso. Si vedrà che come conseguenza di ciò tutti i calcolatori numerici sono in un certo senso equivalenti.

Possiamo ora considerare di nuovo la questione che si era sollevata alla fine del paragrafo 3. Era stato proposto, a titolo di prova, che la domanda "possono pensare le macchine?" venisse sostituita dall'altra "sono immaginabili calcolatori numerici che si comporterebbero bene nel giuoco della imitazione?" Se lo desideriamo possiamo rendere la domanda

ancora più generica e chiedere: "esistono macchine a stati discreti che si comporterebbero bene?" Ma considerando la proprietà della universalità vediamo che entrambe queste questioni sono equivalenti all'ulteriore: "Fissiamo la nostra attenzione su un particolare calcolatore numerico C. È vero che, modificando il calcolatore in maniera da avere a disposizione una memoria adeguata, incrementando adeguatamente la sua velocità di azione e fornendogli una programmazione adeguata, C può prendere soddisfacentemente la parte di A nel giuoco dell'imitazione, se la parte di B viene assunta da un uomo?"

6. Opinioni contrarie a proposito dell'argomento principale

Possiamo ritenere adesso di aver sgombrato il terreno e di essere pronti per cominciare ad esaminare la nostra domanda "possono pensare le macchine?" e la variante ed essa che abbiamo avanzato alla fine dell'ultimo paragrafo. Non possiamo addirittura abbandonare la forma originale del problema, dato che le opinioni sulla legittimità della sostituzione saranno diverse e dobbiamo come minimo tenere presenti quelle che potrebbero essere le obiezioni a questo riguardo. Sarà più semplice per il lettore che io spieghi in primo luogo le mie opinioni in materia. Consideriamo per prima la forma più precisa della domanda. Credo che entro circa 50 anni sarà possibile programmare calcolatori con una capacità di memorizzazione di circa 10^9 , per fare giocare loro il giuoco dell'imitazione così bene che un esaminatore medio non avrà più del 70 per cento di probabilità di compiere l'identificazione esatta dopo 5 minuti di interrogazione. Credo che la domanda iniziale, "possono pensare le macchine?" sia troppo priva di senso per meritare una discussione.

Ciò nonostante credo che alla fine del secolo l'uso delle parole e l'opinione corrente si saranno talmente mutate che chiunque potrà parlare di macchine pensanti senza aspettarsi di essere contraddetto. Credo inoltre che non vi sia alcuna utilità a nascondere queste opinioni. L'opi-

nione popolare che gli scienziati procedano inesorabilmente da un fatto ben stabilito ad un altro fatto ben stabilito, senza che intervenga mai l'influenza di una congettura non ancora provata, è del tutto errata. Purché venga chiaramente messo in evidenza quali siano i fatti provati e quali siano le congetture, non può risultarne alcun danno. Le congetture sono di importanza fondamentale, dato che suggeriscono utili linee di ricerca.

Passo adesso a considerare opinioni opposte alle mie.

L'OBIEZIONE TEOLOGICA. "Il pensare è una funzione dell'anima immortale dell'uomo. Dio ha dato un'anima immortale ad ogni uomo e donna, ma non agli altri animali o alle macchine. Perciò nessun animale o macchina può pensare".

Sono incapace di accettare qualsiasi parte di questa affermazione, ma cercherò di rispondere in termini teologici. Troverei l'argomento più convincente se gli animali fossero classificati insieme agli uomini, perchè c'è una maggior differenza, a mio avviso, tra il tipico animato e l'inanimato, di quanta non ve ne sia tra l'uomo e gli altri animali. Il carattere arbitrario del punto di vista ortodosso diviene più evidente se consideriamo come esso potrebbe apparire ai membri di una diversa comunità religiosa. Come considerano i Cristiani il punto di vista musulmano secondo il quale le donne non hanno anima? Ma lasciamo da parte questo punto e torniamo all'argomento principale. Mi sembra che l'argomento sopra indicato implichi una seria restrizione della onnipotenza di Dio⁽¹⁾. È ammesso che vi siano certe cose che Egli non può fare, come per esempio rendere uno eguale a due, ma non dobbiamo credere che Egli abbia la libertà di concedere l'anima ad un elefante se lo considera opportuno? Potremmo

⁽¹⁾ Può darsi che questa opinione sia eretica. San Tommaso d'Aquino (*Summa theologiae*, citata da Bertrand Russell, *Storia della filosofia occidentale*) dice che Dio non può far sì che un uomo non abbia anima. Ma questa può non essere una effettiva restrizione dei Suoi poteri, ma solo un risultato del fatto che l'anima degli uomini è immortale, e quindi indistruttibile.

aspettarci che Egli eserciterebbe questo potere soltanto in relazione ad una mutazione che fornisse all'elefante un cervello opportunamente migliorato per venire incontro ai bisogni di quest'anima. Si può argomentare con forma esattamente analoga nel caso delle macchine. L'argomento può sembrare diverso perché più difficile da "mandar giù". Ma in effetti questo significa soltanto che noi pensiamo che sarebbe meno probabile che Egli considerasse le circostanze come adatte per conferire un'anima. Le circostanze in questione sono discusse nel resto di questo articolo. Tentando di costruire macchine del genere non vogliamo usurpare irriverentemente il Suo potere di creare anime, più di quanto non facciamo procreando bambini: piuttosto noi siamo, in entrambi i casi, strumenti del suo volere, fornendo abitazioni per le anime che Egli crea.

Comunque, questa è nera speculazione. Non rimango molto impressionato dagli argomenti teologici, qualunque tesi essi sostengano. Argomenti del genere sono risultati insoddisfacenti in passato. Al tempo di Galileo si pensava che i testi, "Ed il sole si fermò... e non si affrettò a tramontare per quasi un giorno intero" (Giosuè 10.13) e "Pose le fondamenta della terra, in modo che non si muovesse per sempre" (Salmi 105.5) costituissero una confutazione adeguata della teoria copernicana. Con le nostre odierne cognizioni un argomento del genere appare futile. Quando queste cognizioni non erano disponibili, esso faceva un'impressione del tutto diversa.

L'OBIEZIONE DELLA "TESTA NELLA SABBIA". "Le conseguenze delle macchine pensanti sarebbero terribili; speriamo e crediamo che esse non possano esistere."

L'argomento è espresso raramente in una forma così esplicita. Ma esso contagia la maggioranza di quanti di noi pensano al problema. Ci piace credere che l'uomo sia in qualche modo misterioso, superiore al resto del creato. Meglio per lui se può dimostrare di essere *necessariamente* superiore, perché allora non vi sarà pericolo che possa perdere la sua po-

sizione di comando. La popolarità dell'argomento teologico è chiaramente connessa con questo sentimento. Probabilmente esso è molto forte tra gli intellettuali, dato che costoro valutano il potere del pensiero più degli altri, e sono più inclini a basare la loro credenza nella superiorità dell'uomo su questo potere.

Non credo che l'argomento sia abbastanza solido per meritare una confutazione. Sarebbe più appropriata una consolazione, che potrebbe magari essere ricercata nella migrazione delle anime.

L'OBIEZIONE MATEMATICA.

Esiste una serie di risultati della logica matematica che possono essere usati per dimostrare che esistono delle limitazioni ai poteri delle macchine a stati discreti. Il più conosciuto di questi è noto come teorema di Gödel, e dimostra che in qualsiasi sistema logico sufficientemente *potente* possono essere formulati degli enunciati che non possono essere né provati né confutati all'interno del sistema, a meno che il sistema stesso non sia contraddittorio. Vi sono altri risultati, sotto qualche aspetto simili, dovuti a Church, Kleene, Rosser e Turing. Il più conveniente da esaminare è quest'ultimo, dato che si riferisce direttamente alle macchine, mentre gli altri possono essere usati soltanto in argomentazioni relativamente indirette: per esempio, se vogliamo utilizzare il teorema di Gödel, abbiamo bisogno di qualche metodo addizionale per descrivere i sistemi logici in termini di macchine, e le macchine in termini di sistemi logici. Il risultato in questione si riferisce ad un tipo di macchina che è essenzialmente un calcolatore numerico a capacità infinita. Esso dice che vi sono alcune cose che una macchina non può fare. Se essa è costretta a dare risposte a domande come nel giuoco della imitazione, ci saranno alcune domande alle quali essa o darà una risposta errata, o non darà affatto risposta, quale che sia il tempo concesso per rispondere. Vi possono naturalmente essere molte domande del genere, e domande cui non può essere data risposta da una macchina possono ricevere una risposta soddisfacente da un'altra macchina.

Stiamo supponendo naturalmente per il momento che le domande siano del tipo per il quale è appropriata una risposta come "sì" e "no", piuttosto che del tipo "Che ne pensi di Picasso?". Le domande alle quali sappiamo che le macchine non riusciranno a rispondere sono di questo tipo: "Si prenda in esame la macchina caratterizzata nella seguente maniera... questa macchina risponderà mai "sì" a qualche domanda?" I puntini devono essere sostituiti dalla descrizione di qualche macchina in una forma standard, descrizione eventualmente analoga a quella del paragrafo 5. Quando la macchina descritta è in una certa relazione relativamente semplice con la macchina che è sottoposta a interrogazione, si può dimostrare che la risposta sarà errata o non sarà data affatto. Questo è il risultato matematico: si sostiene che esso dimostra una incapacità della macchina alla quale l'intelletto umano non è soggetto.

La risposta più breve a questa argomentazione è che sebbene sia stato chiarito che esistono delle limitazioni ai poteri di una qualsiasi macchina specifica, è stato poi soltanto enunciato, senza alcuna sorta di dimostrazione, che nessuna limitazione del tipo è applicabile all'intelletto umano. Non credo però che l'argomentazione possa venire liquidata così sbrigativamente. Ogni volta che viene posta ad una di queste macchine un'opportuna domanda critica, ed essa dà una risposta definitiva, noi sappiamo che questa risposta deve essere errata, e questo ci dà un certo senso di superiorità. Questo senso di superiorità è illusorio? Certamente esso è genuino, ma non credo che vi si possa attribuire troppa importanza. Diamo troppo spesso risposte errate anche noi, per sentirci giustificati nel provar piacere dinanzi a tali prove della possibilità di errore da parte della macchina. Per giunta, la nostra superiorità può essere sentita volta a volta in relazione alla specifica macchina sulla quale abbiamo riportato il nostro piccolo trionfo. Non ci sarebbe alcuna possibilità di trionfo simultaneo su *tutte* le macchine. In breve, quindi, ci possono essere uomini più abili di una qualsiasi macchina data, ma ci

possono poi essere altre macchine più abili ancora, e così via.

Coloro che si rifanno all'argomento matematico sarebbero disposti, credo in maggioranza, ad accettare il giuoco dell'imitazione come base di discussione. Coloro che ricorrono alle due obiezioni precedenti non credo, invece, siano disposti ad accettare alcun criterio.

L'ARGOMENTO DELL'AUTO-COSCIENZA. Questo argomento fu espresso molto bene nel 1949 dal professor Jefferson: "Fino a quando una macchina non potrà scrivere un sonetto o comporre un concerto in base a pensieri ed emozioni provate, e non per la giustapposizione casuale di simboli, non potremo essere d'accordo sul fatto che una macchina eguagli il cervello - cioè, che non solo scriva ma sappia di aver scritto. Nessun meccanismo potrebbe sentire (e non semplicemente segnalare artificialmente, ché sarebbe un facile trucco) piacere per i suoi successi, dolore quando una sua valvola fonde, arrossire per l'adulazione, sentirsi depresso per i propri errori, essere attratto dal sesso, arrabbiarsi o abbattersi quando non può ottenere quel che desidera".

Questo argomento sembra una negazione della validità del nostro esperimento. Secondo la forma più estrema di questa opinione il solo modo per cui si potrebbe essere sicuri che una macchina pensa è quello di *essere* la macchina e di sentire sé stessi pensare. Uno potrebbe allora naturalmente descrivere queste sensazioni al mondo, ma ovviamente nessuno sarebbe giustificato nel darvi ascolto. Allo stesso modo, secondo questa opinione la sola via per sapere che *un uomo* pensa è quella di essere quell'uomo in particolare. È questo in effetti il punto di vista solipsistico. Può essere il punto di vista migliore cui attenersi sul piano logico, ma rende difficile la comunicazione delle idee. Probabilmente A crederà "A pensa, ma B no", mentre B crede "B pensa, ma A no". Invece di discutere in continuazione su questo punto, è normale attenersi alla educata convinzione che ognuno pensi.

Sono sicuro che il professor Jefferson non desidera adottare il punto di

vista estremista e solipsista. Probabilmente egli sarebbe volentieri disposto ad accettare il giuoco dell'imitazione come prova. Il giuoco (con l'esclusione del giocatore B) è frequentemente usato in pratica sotto il nome di esame orale per scoprire se qualcuno ha realmente capito qualcosa o l'ha imparata *a pappagallos*. Ascoltiamo parte di un tale esame orale:

Esaminatore: Nel primo verso del sonetto, che dice "Ti paragonerò a una giornata d'estate", "una giornata di primavera" non andrebbe bene lo stesso?

Candidato: Non quadrerebbe metricamente.

Esaminatore: E "una giornata d'inverno"? Metricamente andrebbe bene.

Candidato: Sì, ma nessuno vorrebbe essere paragonato ad un giorno d'inverno.

Esaminatore: Lei direbbe che Mr. Pickwick le ricorda Natale?

Candidato: In un certo senso.

Esaminatore: Eppure Natale è un giorno d'inverno, e non credo che il paragone dispiacerebbe a Mr. Pickwick.

Candidato: Non credo che lei parli seriamente. Per "un giorno d'inverno" si intende un tipico giorno d'inverno, piuttosto che un giorno speciale come Natale.

E così via.

Che cosa direbbe il professor Jefferson se la macchina che scrive sonetti potesse rispondere in maniera analoga ad un esame orale? Non so se considererebbe la macchina "in fase di segnalazione puramente artificiale" di queste risposte, ma se le risposte fossero così soddisfacenti e ragionate come nel passo indicato sopra, non credo che le descriverebbe come "un facile espediente". Questa frase ritengo sia intesa a tener conto di trucchi come l'inclusione nella macchina di un disco di qualcuno che legga un sonetto, con interruttori adatti per metterlo in moto di volta in volta.

In breve, credo che la maggior parte di coloro che sostengono l'argomento dell'autocoscienza potrebbero es-

sere persuasi ad abbandonarlo, piuttosto che accettare la posizione solipsistica. Quindi essi sarebbero probabilmente disposti a riconoscere la validità della nostra prova.

Non voglio dare l'impressione di credere che non ci sia alcun mistero nei riguardi dell'autocoscienza. C'è, per esempio, qualcosa di paradossale in ogni tentativo di localizzarla. Ma non credo che questi misteri debbano necessariamente essere risolti prima che noi possiamo rispondere alle domande contenute in questo articolo.

ARGOMENTI FONDATI SU INCAPACITÀ VARIE Questi argomenti prendono la forma, "vi concedo che possiate far fare alle macchine tutto quello cui avete accennato, ma non potrete mai costruirne una capace di fare X". Vengono proposte a questo riguardo numerose caratteristiche. Ecco una scelta: Essere gentile, pieno di risorse, bello, cordiale, avere iniziativa, avere senso dell'humour, distinguere il bene dal male, commettere errori, innamorarsi, gustare le fragole con la panna, far sì che qualcuno si innamori di noi, imparare dall'esperienza, usare le parole nel modo appropriato, essere l'oggetto dei propri pensieri, avere un comportamento vario quanto quello umano, fare qualcosa di realmente nuovo.

In genere non viene offerta alcuna base per queste affermazioni. Credo che per la maggior parte siano fondate sul principio della induzione scientifica. Un uomo ha visto durante la propria vita migliaia di macchine. Da quanto ha visto egli trae una serie di conclusioni generali. Esse sono brutte, ognuna è progettata per uno scopo molto limitato, quando vengono usate per uno scopo leggermente diverso si dimostrano inutili, la varietà di comportamento di ognuna è molto limitata, ecc. Naturalmente egli conclude che queste sono proprietà necessarie delle macchine in generale. Molte di queste limitazioni sono associate con la piccolissima capacità di memorizzazione della maggior parte delle macchine. (Sto supponendo che l'idea di capacità di memorizzazione sia estesa in qualche modo a includere macchine diverse da quelle a stati discre-

ti. La definizione esatta non ha importanza, visto che nell'attuale discussione non si richiede alcuna precisione matematica). Pochi anni fa, quando si parlava ancora poco di calcolatori numerici, era possibile suscitare molta incredulità riguardo ad essi, se si accennava alle loro proprietà senza descrivere il modo in cui erano costruiti. Ciò era presumibilmente dovuto ad un'analoga applicazione del principio d'induzione scientifica. Queste applicazioni del principio sono naturalmente in gran parte a livello inconscio. Quando un bambino che si è scottato ha paura del fuoco e mostra di averne paura evitandolo, direi che sta applicando l'induzione scientifica. (Potrei naturalmente descrivere il suo comportamento in molti altri modi).

Le opere e i costumi dell'uomo non sembrano costituire un buon terreno d'applicazione per l'induzione scientifica. Bisogna prendere in esame un grande intervallo nello spazio e nel tempo, se si vogliono ottenere risultati attendibili. Altrimenti potremmo decidere (come fa la maggior parte dei bambini inglesi) che tutti parlano inglese e che è stupido imparare il francese. È necessario tuttavia svolgere alcune considerazioni specifiche a proposito delle incapacità che abbiamo citato. L'incapacità a gustare le fragole con la panna può aver colpito il lettore come cosa frivola. Probabilmente si potrebbe trovare il modo di far gustare ad una macchina questo piatto delizioso, ma ogni sforzo del genere sarebbe sciocco. Ciò che è importante riguardo a questa incapacità è il fatto che essa è connessa ad altre incapacità, ad esempio alla difficoltà che si stabilisca tra l'uomo e la macchina un rapporto di amicizia simile a quello tra bianco e bianco o tra negro e negro. La pretesa che "le macchine non possono sbagliare" sembra strana. Si è tentati di ribattere: "Sono forse peggiori per questo?" Ma cerchiamo di assumere un atteggiamento più comprensivo e di renderci conto del suo vero significato. Penso che questa critica possa venir spiegata in termini di giuoco dell'imitazione. Si afferma che colui che interroga potrebbe distinguere la macchina dall'uomo semplice-

mente ponendo ad entrambi un certo numero di problemi aritmetici. La macchina verrebbe smascherata per la sua tremenda precisione. La risposta a questo è semplice. La macchina, programmata per giocare il giuoco, non cercherebbe di dare la risposta *esatta* a problemi aritmetici. Introdurrebbe deliberatamente degli errori, in un modo studiato apposta per confondere chi interroga. Un difetto meccanico si tradurrebbe probabilmente in una decisione inopportuna sul tipo d'errore da commettere in aritmetica. Anche questa interpretazione della critica non è abbastanza comprensiva. Ma non possiamo permetterci di dedicare ad essa molto spazio. Mi sembra che questa critica dipenda da una confusione tra due tipi di errore. Possiamo chiamarli "errori di funzionamento" e "errori di conclusione". Gli errori di funzionamento sono dovuti a qualche difetto meccanico od elettrico che determina un comportamento della macchina diverso da quello che era stato previsto per essa. In una discussione filosofica si preferisce ignorare la possibilità di errori simili; si discute quindi di "macchine astratte". Queste macchine astratte sono finzioni matematiche più che oggetti fisici. Sono incapaci, per definizione, di errori di funzionamento. In questo senso possiamo veramente dire che "le macchine non possono mai sbagliare". Errori di conclusione possono verificarsi soltanto quando si attribuisce un certo significato ai segnali in uscita della macchina. La macchina potrebbe, ad esempio, presentare equazioni matematiche, o frasi in inglese. Quando viene presentata una proporzione falsa diciamo che la macchina ha commesso un errore di conclusione. Non c'è ovviamente alcun motivo per dire che una macchina non può commettere un errore di questo tipo. Potrebbe non fare niente altro che presentare a ripetizione $0 = 1$. Per fare un esempio meno cattivo, potrebbe avere qualche metodo per tirare conclusioni secondo l'induzione scientifica. Dobbiamo aspettarci che questo metodo possa talvolta condurla a risultati erronei.

All'affermazione che una macchina non può essere oggetto del proprio pensiero si può naturalmente rispondere soltanto se si riesce a dimostrare che la macchina ha un *qualche* pensiero su un *qualche* oggetto. Nondimeno l'espressione "oggetto delle operazioni di una macchina" sembra avere un certo significato, almeno per le persone che si occupano di tali operazioni. Se, per esempio, la macchina tentava di trovare una soluzione all'equazione $X^2 - 40X - 11 = 0$, si sarebbe tentati di descrivere quest'equazione come parte dell'argomento di cui si occupava la macchina in quel momento. In questo senso specifico una macchina può indubbiamente occuparsi di sé stessa. Può essere usata per contribuire ad elaborare i propri programmi o a predire gli effetti di mutamenti nella propria struttura. Osservando i risultati del proprio comportamento può modificare i suoi programmi in modo da conseguire con maggior efficacia un certo scopo. Queste sono possibilità dell'immediato futuro, piuttosto che sogni da utopia.

La critica che una macchina non può presentare molta varietà di comportamento è solo un modo per dire che non può avere molta capacità di memorizzazione. Fino a pochissimo tempo fa una capacità di memorizzazione anche di mille unità era rarissima.

Le critiche che stiamo considerando qui sono spesso forme dissimulate dell'argomento dell'autocoscienza. Normalmente se uno sostiene che una macchina *può* fare una di queste cose, e descrive il tipo di metodo che la macchina potrebbe usare, non fa molta impressione. Si pensa che il metodo (qualunque sia, dato che deve essere meccanico) è in realtà piuttosto banale. (Confronta la parentesi nella frase di Jefferson citata in precedenza).

L'OBIEZIONE DI LADY LOVE-LACE. Le informazioni più dettagliate che possediamo sulla macchina analitica di Babbage sono tratte da un saggio di Lady Lovelace (v. bibliografia). In esso si afferma: "La macchina analitica non ha la pretesa di *creare* alcunché. Può fare *qualsia-*

si cosa sappiamo come ordinarle di fare" (il corsivo è nel testo). Questa affermazione è citata da Hartree, il quale aggiunge: "questo non implica la impossibilità di costruire un'apparecchiatura elettronica capace di "pensare per proprio conto", o nella quale, in termini biologici, si potesse instaurare un riflesso condizionato, che servisse come base per l'"apprendimento". Se ciò sia possibile o no in linea di principio è una domanda stimolante ed entusiasmante, suggerita da alcuni sviluppi recenti. Ma non sembrava che le macchine costruite o progettate a quel tempo avessero questa proprietà".

Sono perfettamente d'accordo con Hartree su questo fatto. Si noterà che egli non afferma che le macchine in questione non avessero questa proprietà, ma piuttosto che gli elementi a disposizione non incoraggiavano Lady Lovelace a crederlo. È perfettamente possibile che le macchine di cui si parla avessero in un certo senso questa proprietà. Supponiamo infatti che qualche macchina a stati discreti abbia questa proprietà. La macchina analitica di Babbage era un calcolatore numerico universale, in modo che, se le sue capacità di memorizzazione e la sua velocità fossero adeguate, si potrebbe, con un'opportuna programmazione, portarla ad imitare la macchina in questione. Probabilmente questo argomento non venne in mente alla Contessa o a Babbage. Comunque essi non avevano il dovere di sostenere tutto ciò che si poteva sostenere.

L'intera questione sarà nuovamente considerata parlando di macchine che apprendono.

Una variante dell'obiezione di Lady Lovelace afferma che una macchina non può "fare mai veramente qualcosa di nuovo". Si può eludere per il momento l'obiezione col detto "Non c'è nulla di nuovo sotto il sole". Chi può essere sicuro che il "lavoro originale" da lui compiuto non sia stato semplicemente la crescita di un seme gettato dall'insegnamento, o la conseguenza dell'aver seguito principi generali bene noti? Una variante migliore dell'obiezione dice che una macchina non può mai

"prenderci alla sprovvista". Questa affermazione è una sfida più diretta e può essere controbattuta direttamente. Le macchine mi prendono alla sprovvista molto frequentemente. Questo di solito dipende in gran parte dal fatto che non faccio calcoli sufficienti a decidere che cosa aspettarmi che facciano, o piuttosto perché, quantunque faccia dei calcoli, li faccio in modo affrettato, disordinato, rischiando di sbagliare. Magari dico a me stesso: "Penso che la tensione qui dovrebbe essere la stessa che lì, comunque supponiamolo". Naturalmente spesso mi sbaglio, e il risultato è una sorpresa per me, perché quando l'esperimento è concluso questi presupposti sono stati dimenticati. Queste ammissioni da parte mia offrono la possibilità per una serie di conferenze che trattino dei miei erronei procedimenti, ma non gettano alcun dubbio sulla mia attendibilità quando testimonio sulle sorprese che vado sperimentando.

Non mi aspetto che questa risposta metta a tacere il mio critico. Egli dirà probabilmente che simili sorprese sono dovute a qualche atto mentale creativo da parte mia e che nessun merito ne deriva in conseguenza alle macchine. Questo ci riporta all'argomento dell'autocoscienza, e lontano dall'idea di sorpresa. È un tipo di argomentazione che dobbiamo considerare esaurita, ma vale forse la pena di osservare che valutare qualcosa come sorprendente richiede sempre un "atto mentale creativo", tanto nel caso che ciò che sorprende sia provocato da un uomo, quanto nel caso che si tratti di un libro, di una macchina o di qualsiasi altra cosa.

L'opinione che le macchine non possano far nascere sorprese è dovuta spesso a un errore cui sono particolarmente soggetti filosofi e matematici. L'errore consiste nel presupporre che appena un fatto si presenta alla mente, tutte le conseguenze di questo fatto saltino fuori simultaneamente. È un presupposto utile in molte circostanze, ma ci si dimentica troppo facilmente che è falso. Una conseguenza naturale di questo modo di agire è che si presuppone che non ci sia alcun merito nella

semplice elaborazione delle conseguenze di dati e principi generali.

L'ARGOMENTAZIONE FONDATA SULLA CONTINUITÀ DEL SISTEMA NERVOSO. Il sistema nervoso non è certo una macchina a stati discreti. Un piccolo errore d'informazione circa la grandezza di un impulso nervoso che colpisce un neurone, può essere importantissimo per quanto riguarda la grandezza dell'impulso in uscita. Si può sostenere che, stando così le cose, non si può pensare di poter imitare il comportamento del sistema nervoso con un sistema a stati discreti. È vero che una macchina a stati discreti dev'essere diversa da una macchina continua. Ma se ci atteniamo alle condizioni del giuoco dell'imitazione, colui che interroga non sarà in grado di trarre alcun vantaggio dalla differenza. Si può rendere più chiara la situazione se consideriamo qualche altra macchina continua più semplice. Un analizzatore differenziale servirà benissimo allo scopo. (Un analizzatore differenziale è un particolare tipo di macchina, non appartenente alla categoria "a stati discreti", usato per determinati calcoli.) Alcune di queste macchine forniscono le loro risposte in forma dattiloscritta, e quindi sono adatte a prender parte al giuoco. Non sarebbe possibile per un calcolatore numerico predire esattamente che risposte verrebbero date dall'analizzatore differenziale a un problema, ma esso sarebbe perfettamente in grado di dare l'esatto tipo di risposta. Per esempio, se gli si chiede di dare il valore di π (circa 3,1416) sarebbe ragionevole scegliere a caso tra i valori 3,12; 3,13; 3,14; 3,15; 3,16 con le probabilità, poniamo, di 0,05; 0,15; 0,55; 0,19; 0,06. In queste circostanze sarebbe molto difficile per colui che interroga distinguere l'analizzatore differenziale dal calcolatore numerico.

L'ARGOMENTAZIONE DEL COMPORTAMENTO SENZA REGOLE RIGIDE. Non è possibile presentare un complesso di regole che descrivano cosa debba fare un uomo in ogni possibile circostanza. Si potrebbe ad esempio avere una regola che stabilisca che ci si deve

fermare quando si vede un semaforo rosso, e si passi invece se se ne vede uno verde, ma cosa succederebbe se per qualche guasto il rosso e il verde apparissero entrambi insieme? Si potrebbe forse decidere che la cosa più sicura è fermarci. Ma più tardi qualche ulteriore difficoltà potrebbe pure nascere da questa decisione. Tentare di fornire regole di condotta tali da includere qualsiasi eventualità, anche quelle che si presentano a causa dei semafori, appare impossibile. Sono d'accordo su tutto questo.

Da ciò si argomenta che non possiamo essere macchine. Proverò a riprodurre l'argomentazione, ma temo che difficilmente riuscirò a non travisarla. Mi sembra che sia pressappoco di questo tipo: "Se ogni uomo avesse un complesso definito di regole di condotta secondo le quali regolare la sua esistenza non sarebbe meglio di una macchina. Ma non esistono tali regole, così gli uomini non possono essere macchine". È chiarissimo che il termine medio non è distribuito. Non penso che la tesi assuma mai veramente questa forma, ma credo che in sostanza sia questo l'argomento usato. Può esserci tuttavia una certa confusione tra "regole di condotta" e "leggi di comportamento" che può far nascere dei dubbi su questo punto. Per "regole di condotta" intendo prescrizioni come "fermati se vedi una luce rossa" in base alle quali si può agire e di cui si può avere coscienza. Per "leggi di comportamento" intendo leggi di natura applicate al corpo di un uomo come "se lo pizzichi, egli si metterà a strillare". Se sostituiamo "leggi di comportamento" che regolano la sua esistenza" al posto di "regole di condotta in base alle quali egli regola la sua esistenza" nell'argomento citato, la difficoltà del termine medio non distribuito non è più insuperabile. Noi crediamo infatti che non soltanto è vero che essere regolati da leggi di comportamento implica essere una specie di macchina (quantunque non necessariamente una macchina a stati discreti), ma che viceversa essere una macchina di questo tipo implica essere regolati da tali leggi. Comunque non possiamo convincerci

così facilmente della mancanza di leggi complete di comportamento, come della mancanza di regole complete di condotta. L'unica maniera che conosciamo per trovare tali leggi è l'osservazione scientifica, e non conosciamo certamente nessuna circostanza nella quale potremmo dire: "Abbiamo cercato abbastanza. Non esistono leggi del genere".

Possiamo decisamente dimostrare che qualsiasi affermazione di questo tipo sarebbe ingiustificata. Supponiamo infatti che potessimo essere sicuri di trovare simili leggi nel caso esistessero. Allora, data una macchina a stati discreti, dovrebbe essere certamente possibile scoprire con l'osservazione abbastanza su di essa da predire il suo comportamento futuro, e questo in un tempo ragionevole, diciamo un migliaio d'anni. Ma non sembra sia questo il caso. Ho impostato per il calcolatore di Manchester un piccolo programma con soltanto 1000 unità di memoria, in base al quale la macchina cui viene fornito un numero a 16 cifre risponde con un altro in due secondi. Sfiderei chiunque ad imparare da queste risposte abbastanza sul programma, da essere in grado di predire qualsiasi risposta a valori non ancora sperimentati.

L'ARGOMENTAZIONE FONDATA SULLA PERCEZIONE ESTRASENSORIALE. Presuppongo che l'idea della percezione extrasensoriale sia familiare al lettore, e così pure il significato dei quattro argomenti in cui essa si articola, cioè telepatia, chiaroveggenza, precognizione, psicocinesi. Questi imbarazzanti fenomeni sembrano contraddire tutte le nostre comuni idee scientifiche. Saremmo molto lieti di poterli mettere in dubbio! Sfortunatamente le prove statistiche, almeno per la telepatia, sono schiacciati. È molto difficile dare una nuova sistemazione alle proprie idee in modo che possano accordarsi con questi nuovi fatti. Una volta che essi sono stati accettati, sembra che non sia poi un gran salto credere negli spiriti e nei fantasmi. L'idea che i nostri corpi si muovano semplicemente secondo le leggi note della fisica e secondo altre leggi non ancora sco-

perte, ma di tipo simile, sarebbe una delle prime a cadere.

Questa argomentazione mi sembra effettivamente molto forte. Si può rispondere che molte teorie scientifiche sembra rimangano utilizzabili in pratica, nonostante siano in conflitto col fenomeno della percezione extrasensoriale, e che in effetti si può andare avanti benissimo anche se ci si dimentica di essa. Questa è una consolazione piuttosto magra, e sorge il timore che il pensiero sia proprio il tipo di fenomeno in cui la percezione extrasensoriale può dimostrarsi particolarmente rilevante.

Un argomento più specifico basato sulla percezione extrasensoriale potrebbe essere il seguente: "Giochiamo il giuoco dell'imitazione, servendoci come testimoni di un uomo che abbia eccellenti doti di ricezione telepatica e di un calcolatore numerico. Chi interroga può porre domande del genere: A che seme appartiene la carta nella mia destra? L'uomo per telepatia o chiaroveggenza dà la risposta esatta centotrenta volte su quattrocento carte. La macchina può solo indovinare a caso, per esempio dando centoquattro volte la risposta esatta, così chi interroga può giungere alla giusta identificazione." Si presenta qui una possibilità interessante. Supponiamo che il calcolatore numerico contenga un generatore di numeri a caso. Allora sarà naturale servirsene per decidere che risposta dare. Ma allora il generatore di numeri a caso sarà soggetto ai poteri psicocinetici di chi interroga. Forse questa psicocinesi potrebbe far sì che la macchina indovini più spesso di quanto ci si potrebbe aspettare da un calcolo probabilistico, cosicché chi interroga potrebbe ancora non essere in grado di giungere all'identificazione esatta. D'altra parte, potrebbe essere in grado di indovinare senza fare alcuna domanda, per chiaroveggenza. Con la percezione extrasensoriale qualsiasi cosa può accadere.

Se si ammette la telepatia, diventa necessario rivedere il nostro esperimento cautelandosi. La situazione può considerarsi analoga a quella che si verificherebbe se colui che interroga parlasse a sé stesso e uno dei contendenti stesse ascoltando con

l'orecchio alla parete. Una "camera a prova di telepatia" in cui mettere i contendenti soddisferebbe tutte le esigenze.

7. Macchine che apprendono

Il lettore si sarà già accorto che non ho alcun argomento molto convincente di carattere positivo per sostenere il mio punto di vista. Se l'avessi avuto non mi sarei certo dedicato con tanta cura a indicare gli errori dei punti di vista opposti al mio. Le prove che ho le esporrò ora.

Torniamo per un attimo all'obiezione di Lady Lovelace, secondo la quale la macchina può fare solo ciò che noi le diciamo di fare. Si potrebbe dire che un uomo può "iniettare" una idea nella macchina, e che essa risponderà in una certa misura per poi ricadere nello stato di riposo, come una corda di pianoforte colpita da un martello. Un altro paragone potrebbe essere costituito da una pila atomica di grandezza inferiore a quella critica: un'idea immessa corrisponderebbe allora a un neutrone che entra nella pila dal di fuori. Ciascun neutrone di questo tipo determinerà un certo disturbo che eventualmente poi scomparirà. Se, tuttavia, la grandezza della pila viene sufficientemente aumentata, il disturbo causato da un neutrone in ingresso molto probabilmente continuerà ad aumentare fino alla distruzione dell'intera pila. Esiste un fenomeno corrispondente per la mente, e ne esiste uno per le macchine? Per la mente umana sembra che ne esista uno. La maggior parte delle menti umane sembra essere "sotto il livello critico", essere equivalente cioè in questa analogia alle pile di grandezza inferiore a quella critica. Un'idea che si presenti ad una mente di questo tipo farà nascere in media meno di un'idea, in risposta. Una parte piuttosto piccola delle menti umane è invece "sopra il livello critico". Un'idea che si presenti ad una di queste menti può far nascere un'intera teoria fatta di idee del secondo ordine, del terzo o ancora più remote. Attenendoci a questa analogia possiamo chiedere: "Si può fare in modo che una macchina sia sopra il livello critico?"

Anche l'analogia delle "pelle della cipolla" può servire. Considerando le funzioni mentali o quelle del cervello troviamo certe operazioni che possiamo spiegare in termini puramente meccanici. Questo, diciamo, non corrisponde alla mente come essa è in realtà: è una specie di pelle che dobbiamo togliere se vogliamo trovare la mente reale. Ma poi in ciò che rimane troviamo un'altra pelle da togliere, e così via. Procedendo a questo modo arriviamo infine alla mente "reale", o solo ad una pelle che non contiene nulla? Nell'ultimo caso tutta la mente sarebbe di tipo meccanico. (Non si tratterebbe comunque di una macchina a stati discreti. Ne abbiamo già discusso.)

Questi due ultimi paragrafi non pretendono di presentare argomenti convincenti. Dovrebbero piuttosto essere descritti come "tirate" tendenti a far nascere una credenza".

L'unico sostegno veramente soddisfacente che può essere fornito per il punto di vista espresso all'inizio del paragrafo 6 si avrà aspettando la fine di questo secolo ed eseguendo l'esperimento descritto. Ma cosa possiamo dire nel frattempo? Che passi dovremmo compiere ora, perché l'esperimento riesca?

Come ho spiegato, il problema è soprattutto di programmazione. Si dovranno compiere progressi anche nella tecnica ma sembra improbabile che essi non risultino adeguati al bisogno. Le valutazioni della capacità di memorizzazione del cervello vanno da 10^{10} a 10^{15} unità binarie. Io sono incline a considerare più esatte le valutazioni più basse e penso che solo una parte piccolissima venga usata per i tipi superiori di pensiero. La maggior parte viene probabilmente usata per ritenere le impressioni visive. Mi sorprenderei se occorressero più di 10^9 unità per giocare in modo soddisfacente il giuoco dell'imitazione, almeno contro un cieco. (Si noti che la capacità dell'*Enciclopedia britannica*, undicesima edizione, è di 2×10^9 unità binarie.) Una capacità di memorizzazione di 10^7 unità sarebbe senz'altro alla nostra portata anche con le attuali tecniche. Non è probabilmente affatto necessario aumentare la velocità di funzionamento delle macchi-

ne. Le parti delle macchine moderne che possono essere considerate analoghe alle cellule nervose funzionano ad una velocità circa mille volte superiore. Ciò dovrebbe fornire un "margine di sicurezza" sufficiente a compensare le perdite di velocità che in vario modo possono determinarsi. Il nostro problema quindi è di trovare il modo di programmare queste macchine per poter giocare il giuoco. Al mio attuale ritmo di lavoro elaboro in un giorno circa mille unità di programma, cosicché circa sessanta operatori, lavorando continuamente per cinquant'anni, potrebbero portare a termine il lavoro, se nulla dovesse venir cestinato. Appare desiderabile un metodo più rapido.

Cercando di imitare una mente umana adulta siamo tenuti a riflettere parecchio sul processo che l'ha condotta allo stato in cui si trova. Possiamo notare qui tre componenti:

- lo stato iniziale della mente, diciamo alla nascita;
- l'educazione cui è stata sottoposta;
- altre esperienze, che non possono venir descritte come educazione, che essa ha vissuto.

Invece di elaborare un programma per la simulazione di una mente adulta, perché non proviamo piuttosto a realizzarne uno che simuli quella di un bambino? Se la macchina fosse poi sottoposta ad un appropriato corso d'istruzione, si otterrebbe un cervello adulto. Presumibilmente il cervello infantile è qualcosa di simile ad un taccuino di quelli che si comprano dai cartolai. Poco meccanismo e una quantità di fogli bianchi (meccanismo e scrittura sono dal nostro punto di vista quasi sinonimi). La nostra speranza è che ci sia così poco meccanismo nel cervello infantile, che qualcosa di analogo possa venir facilmente programmato. Per il processo educativo possiamo supporre che il lavoro, in prima approssimazione, sia pressappoco uguale a quello necessario per il bambino. Abbiamo così diviso il nostro problema in due parti. Il programma a livello infantile e il processo educativo. Essi sono strettamente connessi. Non possiamo

aspettarci di trovare una buona macchina al primo tentativo. Bisogna sperimentare un metodo d'insegnamento per una macchina del genere e vedere in che misura essa impari. Si può poi provarne un'altra e vedere se è migliore o peggiore. C'è una connessione evidente tra questo processo e l'evoluzione:

Struttura della macchina-bambino → Materiale ereditario

Cambiamenti della macchina-bambino → Mutazioni

Giudizio dello sperimentatore → Selezione naturale.

È da sperare, peraltro, che questo processo sia più rapido dell'evoluzione. La sopravvivenza del più adatto è un metodo troppo lento per quanto concerne la evidenziazione del progresso realizzato. Lo sperimentatore, servendosi della sua intelligenza, dovrebbe essere in grado di renderlo più veloce. Ugualmente importante è il fatto che egli può non limitarsi ad attendere le mutazioni casuali. Se è in grado di scoprire la causa di qualche difetto può probabilmente pensare al tipo di mutazione che lo eliminerebbe.

Non sarà possibile applicare alla macchina proprio lo stesso metodo d'insegnamento che si usa per un bambino. Per esempio, essa non avrà gambe e quindi non le si potrà chiedere di uscire per riempire il secchio del carbone. Potrebbe non avere occhi. Ma anche se questi difetti potessero venir brillantemente superati da abili espedienti meccanici, non si potrebbe mandare una simile creatura a scuola senza farla beffeggiare in modo eccessivo dagli altri bambini. Bisogna le sia data una certa protezione. Non dobbiamo preoccuparci troppo per le gambe, gli occhi, ecc. L'esempio di Helen Keller mostra che il processo educativo si può svolgere, purché vi sia un mezzo per la comunicazione in ambedue le direzioni tra maestro e allievo.⁽²⁾

Normalmente associamo punizioni e ricompense al processo d'insegnamento. Alcune semplici macchine-bambino possono essere costruite o programmate in base a questo tipo di principio. La macchina deve essere costruita in modo che sia impossi-

bile che si ripetano gli avvenimenti che precedettero di poco il verificarsi di un segnale di punizione, mentre un segnale di ricompensa aumenta la probabilità di ripetizione degli avvenimenti che hanno condotto ad esso. Queste definizioni non presuppongono alcun sentimento da parte della macchina. Ho fatto alcuni esperimenti con una simile macchina-bambino e sono riuscito ad insegnarle alcune cose, ma il metodo d'insegnamento era troppo poco ortodosso perché gli esperimenti potessero venir considerati veramente riusciti.

L'uso di punizioni e ricompense può al più costituire una parte del processo d'insegnamento. Parlando approssimativamente, se l'insegnante non ha altri mezzi per comunicare con l'allievo, la quantità di informazione che può fargli giungere non supera il numero totale delle ricompense e delle punizioni assegnategli. Quando un bambino avesse appreso a ripetere "Casabianca", si sentirebbe probabilmente davvero indolenzito, se la parola potesse venire individuata solo con una tecnica del tipo "venti domande", ed ogni "no" si fosse tradotto in una botta. È necessario perciò disporre di altri canali di comunicazione, "non emozionali". Se vi sono questi canali, è possibile mediante punizioni e premi insegnare a una macchina ad eseguire ordini dati in un certo linguaggio, per esempio un linguaggio simbolico. L'uso di questo linguaggio diminuirà notevolmente il numero delle punizioni e delle ricompense necessario.

Ci possono essere vari punti di vista sulla complessità più opportuna per la macchina-bambino. Si potrebbe tentare di farla più semplice possibile purché in accordo coi principi generali. Oppure si potrebbe "incorporarvi" un sistema completo di inferenza logica. In quest'ultimo caso la memoria sarebbe in gran parte occupata da definizioni e proposizioni. Le proposizioni sarebbero di vari tipi, per esempio fatti ben fondati, congetture, teoremi dimostrati matematicamente, dichiarazioni fornite da un'autorità, espressioni aventi la forma logica di una proposizione, ma nessun valore di credibilità.

Alcune proposizioni possono essere descritte come "imperativi". La macchina dovrebbe venir costruita in modo che appena un imperativo viene classificato come "ben fondato" si svolge automaticamente l'azione opportuna. Per illustrare questo, supponiamo che l'insegnante dica alla macchina: "Fa' i tuoi compiti ora." Ciò potrebbe determinare l'inclusione tra i fatti ben fondati dell'espressione "L'insegnante dice: 'Fa' i tuoi compiti ora'." Un altro fatto simile potrebbe essere "Ogni cosa che dice l'insegnante è vera". La combinazione di queste espressioni potrebbe eventualmente condurre ad includere tra i fatti ben fondati l'imperativo "Fa' i tuoi compiti ora", e ciò, per il modo in cui la macchina è costruita, significa che essa comincerà effettivamente a fare i compiti, e che il risultato sarà molto soddisfacente. Il processo di inferenza usato dalla macchina non ha bisogno di possedere requisiti tali da soddisfare i logici più esigenti. Potrebbe, per esempio, non esserci alcuna gerarchia di tipi. Ma questo non significa necessariamente che verranno commessi errori di tipo, quant'è vero che non siamo obbligati a cadere da un dirupo per il fatto che non è recintato. Imperativi opportuni (espressi nell'ambito dei sistemi, non compresi nelle regole del sistema) come "Non usare una classe a meno che non sia una sottoclasse di una di cui ha parlato l'insegnante" possono avere un effetto simile a "Non avvicinarti troppo all'orlo".

Gli imperativi cui può obbedire una macchina che non ha membra sono necessariamente di tipo piuttosto intellettuale, come nell'esempio fornito sopra (fare i compiti). Avranno speciale importanza tra questi impe-

(2) Scrittrice americana (n. 1880), famosa soprattutto per la straordinaria vicenda della sua educazione in condizioni di gravissima minorazione fisica. Cieca, sorda e muta dai due anni, all'età di sei fu affidata a una ex cieca parzialmente guarita, Ann Mansfield Sullivan, allora ventenne. Le eccezionali doti di insegnante della Sullivan (la *Anna dei miracoli* del noto dramma e del film ricavatone) e l'acuta sensibilità e intelligenza della Keller ebbero come risultato che la bambina non solo capì cosa sia il linguaggio, ma giunse a poter frequentare scuole speciali e addirittura a conseguire il diploma *cum laude* di un istituto d'istruzione superiore.

rativi quelli che regolano l'ordine in cui devono essere applicate le regole del sistema logico in questione. Ad ogni livello infatti, quando si usa un sistema logico, c'è un grandissimo numero di passi diversi, ciascuno dei quali può essere compiuto, pur rispettando le regole del sistema logico in questione. Queste scelte mettono in luce la differenza tra un ragionatore brillante ed uno sciocco, non la differenza tra un ragionamento valido e uno errato. Le proposizioni che portano a imperativi di questo tipo potrebbero essere "quando si cita Socrate usare il sillogismo in *barbara*"⁽³⁾ oppure "Se è stato provato che un metodo è più rapido di un altro, non usare il metodo più lento". Qualcuna di esse può essere fornita alla macchina "d'autorità", ma altre possono essere prodotte dalla macchina stessa, per esempio per induzione scientifica. L'idea di una macchina che impara può apparire paradossale ad alcuni lettori. Come possono cambiare le regole di funzionamento della macchina? Esse dovrebbero descrivere completamente come reagirà la macchina qualsiasi possa essere la sua storia, a qualsiasi cambiamento possa essere soggetta. Le regole sono quindi assolutamente invarianti rispetto al tempo. Questo è verissimo. La spiegazione del paradosso è che le regole che vengono cambiate nel processo di apprendimento sono di tipo meno pretenzioso e intendono avere solo una validità temporanea. Il lettore può fare un parallelo con la costituzione degli Stati Uniti.

Una caratteristica importante di una macchina che impara è che il suo insegnante ignorerà spesso in gran parte ciò che di preciso si verifica al suo interno, quantunque possa essere in grado di predire in qualche misura il comportamento del suo allievo. Questo dovrebbe valere nel modo più deciso per l'educazione successiva di una macchina nata da una macchina-bambino ben progettata (o programmata). Ciò contrasta nettamente con la procedura normale quando si usa una macchina per fare calcoli: lo scopo è allora di avere una chiara immagine mentale dello stato della macchina in ogni momento del calcolo. Questo scopo può esse-

re conseguito con uno sforzo. La teoria che "la macchina può fare soltanto ciò che sappiamo come ordinarle"⁽⁴⁾, sembra strana se consideriamo tutto questo. La maggior parte dei programmi che possiamo inserire nella macchina avranno come risultato di farle fare qualcosa che non possiamo assolutamente capire, o che giudichiamo come comportamento completamente casuale. Il comportamento intelligente consiste presumibilmente nello staccarsi dal comportamento completamente prevedibile implicato nel calcolo, ma di poco, in modo da non determinare un comportamento casuale o dei giri viziosi che si risolvono in inutili ripetizioni. Preparando la macchina per la sua parte nel giuoco dell'imitazione mediante un processo di insegnamento e di apprendimento, conseguiamo un altro importante risultato: è probabile cioè non occorra più preoccuparsi della "capacità dell'uomo di commettere errori" imitandola in modo naturale, ossia senza un "tirocinio speciale". I processi che si apprendono non assicurano al cento per cento il risultato: altrimenti non potrebbero essere disimparati.

È probabilmente un buon provvedimento introdurre un elemento casuale in una macchina che impara. Un elemento casuale è piuttosto utile quando cerchiamo la soluzione di qualche problema. Supponiamo, per esempio, di voler trovare un numero tra 50 e 200 uguale al quadrato della somma delle sue unità; potremmo cominciare con 51, poi con 52 e continuare fino ad ottenere un numero che vada bene. Come alternativa, potremmo scegliere numeri a caso fino a trovarne uno buono. Tale metodo ha il vantaggio che non è necessario tenere conto dei valori già provati, ma lo svantaggio che sussiste la possibilità di provare due volte lo stesso numero; questo però non ha molta importanza se esistono diverse soluzioni. Il metodo sistematico ha lo svantaggio che può esserci un blocco enorme senza alcuna soluzione proprio nel gruppo che deve essere esaminato per primo. Ora il processo di apprendimento può essere considerato come la ricerca di una forma di comporta-

mento che soddisferà l'insegnante (o qualche altro criterio). Dato che esiste probabilmente un gran numero di soluzioni soddisfacenti, il metodo casuale sembra migliore di quello sistematico. Bisogna notare che esso è impiegato nell'analogo processo dell'evoluzione. Ma in quel caso il metodo sistematico non è possibile. Come si potrebbe tener conto delle diverse combinazioni genetiche che sono state tentate, in modo da evitare di sperimentarle di nuovo?

Possiamo sperare che le macchine saranno alla fine in grado di competere con gli uomini in tutti i campi puramente intellettuali. Ma quali sono i migliori per cominciare? Anche questa è una decisione difficile. Molta gente pensa che un'attività molto astratta, come giocare a scacchi, sarebbe la migliore. Si può anche sostenere che è meglio fornire alla macchina i migliori organi di senso che si possano comprare e poi insegnarle a capire e parlare l'inglese. Questo processo potrebbe seguire il metodo d'insegnamento normale per un bambino. Le cose verrebbero indicate, verrebbe dato loro un nome, ecc. Ancora una volta ignoro quale sia la risposta esatta, ma penso che bisognerebbe tentare ambedue le strade. Possiamo vedere nel futuro solo per un piccolo tratto, ma possiamo pure vedere che in questo piccolo tratto c'è molto da fare.

Bibliografia

- [1] S. BUTLER, *Erewhon* (Londra 1865) trad. L. Drudi Demby (Adelphi, Milano 1965).
- [2] A. CHURCH, *An Unsolvable Problem of Elementary Number Theory*, Amer. J. Math., vol. 58, 345-63 (1936).
- [3] K. GÖDEL, *Über formal unentscheidbare Sätze der "Principia mathematica" und verwandter Systeme, I*, Mh. Math. Phys., vol. 38, 173-89 (1931); trad. ital. in E. Agazzi, *Introduzione ai problemi dell'assiomatica* (Milano 1961).

(3) Nome convenzionale dato dagli scolastici al sillogismo della forma "Se ogni B è A, e ogni C è B, allora ogni C è A. Aristotele ne parla negli *Analitici primi*.

(4) Confronta l'affermazione di Lady Lovelace in cui non compare la parola "soltanto".

[4] D.R. HARTREE, *Calculating Instruments and Machines* (New York 1949).

[5] S.C. KLEENE, *General Recursive Functions of Natural Number*, Amer. J. Math., vol. 57, 153-73, 219-44 (1935).

[6] G. JEFFERSON, *The Mind of Mechanical Man* (Lister Oration for 1949), Brit. med. J., vol. 1, 1105-21 (1949).

[7] LADY LOVELACE, *Translator's Notes to an Article on Babbage's Analytical Engine*, nelle "Scientific Memoirs" a cura di R. Taylor, vol. 3 (1842) pp. 691-731. Tali note sono state integralmente ristampate nell'appendice 1 del volume *Faster than Thought: A Symposium on Digital Machines*, a cura di B.V. Bowden (Pitman, Londra 1953).

[8] B. RUSSEL, *Storia della filosofia occidentale*, trad. L. Pavolini (Longanesi, Milano 1963).

[9] A.M. TURING, *On Computable Numbers, with an Application to the Entscheidungsproblem*, Proc. Lond. math. Soc., (2) vol. 42, 230-65 (1936).

La produzione del software applicativo: evoluzione e prospettive

a cura di Enrico Spoletini

È uscito il numero 10 della "Collana dei Quaderni di Informatica" edita da Franco Angeli e curata dalla Honeywell Information Systems Italia. Si tratta di **"La produzione del software applicativo: evoluzione e prospettive"**, di vari Autori, a cura di Enrico Spoletini.

L'applicazione dell'informatica trova sempre più i suoi vincoli non tanto nelle caratteristiche economiche o funzionali dell'hardware quanto nei costi e nei tempi di sviluppo del software. È sufficiente ricordare che oggi un microprocessore costa meno di una riga di buon software applicativo.

Questo stato di cose riflette l'andamento asimmetrico della evoluzione tecnica dei due settori. A fronte dell'incessante, radicale innovazione delle basi concettuali e metodologiche dell'hardware, il mutamento delle tecniche di fabbricazione del software è stato in proporzione modesto.

I livelli di automazione, standardizzazione e modularità sono completamente diversi nei due campi. Basti pensare alla diversa complessità dei «mattoni» costruttivi: laddove nell'hardware si lavora con componenti integrati, migliaia di volte più complessi dei dispositivi di un tempo, nel software i componenti sono rimasti sostanzialmente al livello primordiale dell'istruzione.

Questo divario tra hardware e software tende ad allargarsi ed a condizionare l'impiego dell'informatica.

Diventa quindi sempre più importante cercare di aumentare la produttività del software.

L'aumento di produttività non può però basarsi su fatti estemporanei o individuali, ma richiede miglioramenti di natura sistematica, applicabili alla media degli operatori. Ecco che ci si riporta quindi al discorso metodologico come punto cruciale per l'evoluzione dell'area del software.

