

Argomenti di psicologia

NUOVE TEORIE DELLA MENTE

**CONCEZIONI RECENTI
SU MENTE, PENSIERO, INTELLIGENZA**

a cura di

Mario Groppo - Alessandro Antonietti



VITA E PENSIERO

Pubblicazioni dell'Università Cattolica

Loredana Cena

L'INTELLIGENZA ARTIFICIALE: UN PROFILO STORICO

1. "Cogito ergo sum" e "computo ergo sum"? Che cosa si intende per Intelligenza Artificiale

Per gli esseri umani il concepirsi viventi e coscienti viene strettamente messo in relazione con la nozione di funzionamento mentale. Il famoso detto di Cartesio "*Cogito, ergo sum*", sarebbe indicativo del fatto che l'uomo trae la prova della sua esistenza proprio dalla consapevolezza dei suoi processi cognitivi.

Ma qual è la natura di questi processi? È possibile rappresentare con una macchina il pensiero, i sentimenti e le emozioni mediante una struttura di regole, in modo tale che possano venire poi riprodotte artificialmente? Se si riuscisse a simulare i procedimenti dell'intelligenza umana con un computer, questi potrebbe ragionare e pensare come un essere umano? Un computer potrebbe un giorno diventare capace di elaborare e comunicarci "*Computo, ergo sum*"?

Una delle fantasie più ricorrenti nell'uomo è sempre stata creare macchine con cui parlare, capaci di pensare da sole e che magari un giorno possano persino sostituirlo. Questa realtà è per ora ancora molto lontana e, secondo il filosofo Hubert Dreyfus, sarà poco plausibile una formalizzazione della nostra comprensione in un sistema complesso di fatti e regole (Dreyfus, 1972). Pur tuttavia questo è uno degli scopi per cui ha preso vita quella scienza denominata Intelligenza Artificiale, che tuttora stimola studiosi di diverse discipline.

Ma che cosa si intende per Intelligenza Artificiale?

L'IA, l'Intelligenza Artificiale (in inglese AI, *Artificial Intelligence*), è una branca dell'Informatica che si prefigge di emulare, mediante sistemi automatici, comportamenti intelligenti, tipici degli esseri umani; essa comprende quella serie di tentativi che gli scienziati stanno compiendo per imitare il funzionamento della mente dell'uomo. Nata

semplice elaborazione dell'informazione: le persone la comprendono, "traggono un senso" da ciò che odono e vedono; hanno nuove idee che apparentemente nascono dal nulla; usano il senso comune per muoversi in un mondo che talvolta appare molto illogico (Rich, 1983). Inoltre, proprio le abilità che ci appaiono come le più automatiche, in cui non sembra essere necessaria molta intelligenza, si sono dimostrate molto difficili da simulare su un computer (Simon, 1973; Bara, 1977, 1978). Questo sorprese i primi ricercatori di IA che ritenevano che le azioni umane più elementari potessero venire programmate facilmente in un computer: si rivelò invece il contrario. Più un compito è complesso per la mente umana, più questa utilizza il pensiero cosciente e volontario. Nella programmazione per computer è necessario determinare ogni singolo passaggio, secondo una successione finalizzata a raggiungere un certo risultato e questa operazione può essere molto complessa. Inoltre, se una azione è così naturale per una persona che non deve nemmeno pensarci nell'eseguirla, si può trovare difficoltà nella descrizione di come è stata compiuta passo per passo.

La ricca rete di interconnessioni nella memoria umana è una delle differenze più profonde fra uomo e macchina. La capacità del cervello di ricercare simultaneamente l'informazione attraverso i suoi milioni di neuroni sembra decisamente fantastica. A differenza del computer, quanta più informazione abbiamo circa un argomento, tanto più rapidamente riusciamo a richiamarla, ma la memoria umana a volte è tutt'altro che perfetta. L'informazione in certi casi va perduta, perché scomponiamo le esperienze in frammenti che immagazziniamo in parti diverse della memoria: scomporre le conoscenze in memoria ci permetterebbe di fare migliori generalizzazioni e predizioni più utili. Può darsi che le imperfezioni umane non siano per nulla casuali: la nostra memoria a breve termine forse ci costringe a usare l'astuzia. Alcuni dei nostri limiti forse sono trucchi della natura per farci scoprire modi più sottili di fare le cose (Minsky, 1985). L'idea del computer che elabora pedestremente i dati, nonostante la sua capacità di eseguire miliardi di calcoli al secondo, è un altro indicatore di come lo studio dell'IA stia cambiando l'immagine che noi abbiamo di ciò che conta nell'intelligenza.

Venticinque anni fa si riteneva che l'integrazione simbolica nel calcolo integrale fosse autentica intelligenza. Ora ci sono programmi che superano di gran lunga le prestazioni umane in questo campo e sembra invece molto più complessa la nostra capacità di tenere una conversazione o capire una scena, solo guardandola. Queste riflessioni hanno condotto Minsky (1980), a chiedersi «Perché la gente pensa che i

computer non pensino? Come mai siamo riusciti a ottenere dai programmi d'intelligenza artificiale prestazioni così adulte, prima di poterli far fare cose puerili?». Egli risponde che forse la ragione è che il pensiero adulto, o il pensiero specialistico, è spesso in qualche maniera più semplice del senso comune d'un bambino. Ciò che comunemente è inteso come senso comune può essere caratterizzato da procedimenti di rappresentazione del sapere molto complessi. L'occhio di un bambino, ad esempio, è in grado di calcolare una distanza, di riconoscere le persone e di evitare collisioni quando gioca in un campo. Il cervello dello stesso bimbo, in base alle informazioni percepite dalla vista, trasmette messaggi agli arti permettendogli di compiere movimenti complessi, ad esempio come un salto. Michael Brady della Oxford University (Brady, 1978) ritiene che sappiamo come riprodurre con il computer, in un modello, gli spostamenti di persone all'interno di un quadro che si modifica, eppure un bambino è capace di muoversi in mezzo a 100 persone e di arrivare al punto che si è prefissato.

Si delinea pertanto un campo alquanto complesso di indagine a cui convergono esperienze molto diverse: da quelle della filosofia a quelle dell'economia, psicologia e cibernetica, pedagogia e linguistica, matematica e logica, sociologia e statistica; entro metodi e finalità comuni lavorano neurofisiologi, ingegneri elettronici ed informatici. Vediamo pertanto di fare un punto sullo stato dell'arte di quella che si presenta come una affascinante area di ricerca.

2. Le macchine potranno mai "pensare"?

Dalla conferenza di Dartmouth alla rivoluzione IA

Prima di passare alla storia è necessario fare alcuni brevi cenni su quella che è la preistoria dell'IA.

Le premesse scientifiche e tecnologiche di questa disciplina affondano le loro radici nell'informatica, con cui essa ha un rapporto di parentela privilegiato. Negli anni che vanno dal 1930 al 1950, alcuni filosofi e matematici, tra cui Russell, Whitehead, Tarsky, Church, Turing, Goedel, contribuiscono a nuove elaborazioni sulla logica, sul ragionamento e sul calcolo. Molte delle loro teorie vengono a costituire le fondamenta dell'attuale ricerca teorica in Intelligenza Artificiale. Negli anni durante la Seconda Guerra Mondiale è la guerra a dare un impulso alle ricerche in questo settore. Per la applicabilità in campo militare vengono finanziati la costruzione del Colossus, nel 1943, che contribuì

proprie cognizioni con un tipo di conoscenza che può essere usata nel raffronto di configurazioni complesse. Questa conversione della conoscenza è forse ciò che è direttamente responsabile dell'efficienza con cui lo specialista esperto applica le sue conoscenze, ma anche della difficoltà che incontra quando si tratta di rendere esplicito il procedimento. L'apprendimento di una lingua straniera può servire da esempio. Per prima cosa si apprendono le regole della lingua e la capacità di parlarla è all'inizio lentissima, perché è necessario tenere continuamente presente nella mente queste regole e capire come si applicano alla situazione del momento. Acquisendo più scioltezza non occorre richiamare le regole, fino a che quasi ce le dimentichiamo. Accedendo sempre meno consciamente al nostro bagaglio di conoscenze, aumentiamo la nostra capacità e scioltezza ad applicarlo. Applicare una conoscenza, però, non è la stessa cosa che acquisirla.

Il problema della rappresentazione della conoscenza e della sua acquisizione è uno dei maggiori nel campo dell'Intelligenza Artificiale. L'apprendimento di competenze formali, come la soluzione di problemi di fisica o la dimostrazione di teoremi di geometria, non avviene semplicemente leggendo un libro di testo e comprendendo i principi astratti, ma risolvendo problemi applicativi: quello che il manuale non insegna è quando applicare le nozioni immagazzinate.

Per quanto riguarda la rappresentazione della conoscenza generalmente si utilizzano due tipi di conoscenza: una conoscenza dichiarativa in cui si conoscono fatti che riguardano oggetti, eventi e relazioni che intercorrono fra di essi e una conoscenza procedurale con cui si procede a utilizzare la conoscenza dichiarativa e si dà un senso. La natura della conoscenza umana non è del tutto nota, e ciò rende più difficile sintetizzarla su un calcolatore: rappresentare la conoscenza in una macchina è stata la questione chiave in tutta la storia dell'Intelligenza Artificiale. Attualmente, non si conosce ancora il metodo "migliore" per rappresentare la conoscenza. Barr e Feigenbaum affermano che «non è possibile trovare uno schema che meglio di altri rappresenti alcuni aspetti della memoria umana». «Non sappiamo perché alcuni schemi si rivelino buoni per alcuni compiti e non per altri. Ma alcuni di essi si sono mostrati molto utili in vari programmi che presentano un comportamento intelligente» (Barr e Feigenbaum, 1982, 147).

Per insegnare al computer occorre trattare il problema della programmazione logica. Essa si serve di due tecniche di base:

- 1) la rappresentazione della conoscenza in maniera logica
- 2) e il suo utilizzo per risolvere i problemi.

Con la programmazione logica il lavoro si sviluppa a ritroso: ciò che può sembrare creativo è in realtà il risultato di un procedimento di analisi sistematica. Diciamo ad esempio:

Fido è un cane dalmata, se Fido è un cane ed è punteggiato di macchie,

oppure diciamo

Fido è un bassotto, se Fido è un cane e se Fido assomiglia ad esempio a una salciccia.

Se quello che ci interessa è la razza riduciamo il problema alle condizioni espresse da certe regole, piuttosto che da altre. Innanzitutto confermiamo che Fido è un cane, poi andiamo a vedere il colore, la forma e a seconda di quale condizione è vera giungiamo alla conclusione: basso o dalmata. Questa operazione è denominata *semantic network*: è una piccola rete semantica in cui si dice che un cane è un tipo di mammifero e fra il cane e il mammifero esiste una relazione. Si possono avere altre relazioni che ci dicono ad esempio che Fido è un cane di una certa razza e che appartiene alla famiglia Rossi. Anche la conoscenza dichiarativa può essere rappresentata graficamente da una rete semantica, in cui i dati sono collegati da interrelazioni. In termini strutturali, una rete semantica è costituita da fatti, uniti da legami che rappresentano le relazioni. In un dominio che comprende una grande quantità di conoscenza con complesse interrogazioni, una rete semantica può costituire le fondamenta di un sofisticato sistema di inferenza.

Ora come usa il computer questa conoscenza?

Il sistema simbolico di ragionamento e di deduzione, codificato dagli antichi greci, quello della logica matematica, era già in uso prima della nascita dei calcolatori. Tale sistema ha dimostrato una semplice applicabilità su calcolatore; i calcolatori sono basati su teorie di logica matematica. La medesima architettura dei calcolatori si conforma alla logica formale, che rappresenta un metodo di rappresentazione della conoscenza in certo modo limitato. Va considerato il fatto che le persone non usano solamente la logica nel proprio comportamento e solo una minima parte delle azioni umane intelligenti possono essere descritte con parametri della logica e della consapevolezza (Imbasciati, 1990).

Un programma di Intelligenza Artificiale si serve della logica per rappresentare la conoscenza e configura il proprio dominio mediante formule. Attraverso l'applicazione di una serie di rigide regole di inferenza alla sua base di conoscenza, il programma può giungere alle conclusioni desiderate. Le reti logiche e semantiche costituiscono siste-

errore, una semplificazione, che riduce i passaggi nello spazio di ricerca per la soluzione di un problema. Non garantisce soluzioni ottime ma fornisce metodi intelligenti per analizzare i problemi, al di là dell'approccio "prova ed errore" di una ricerca cieca (Feigenbaum e Feldman, 1963). La scelta di una tecnica di ricerca nell'IA, dato che non esiste una modalità di rappresentare la conoscenza universalmente accettata come la più valida, dipende dal problema, per cui può a seconda dei casi risultare più opportuno condurre una ricerca di tipo orizzontale (*breadth first*) o di tipo verticale (*depth first*) oppure utilizzare un concatenamento diretto (*forward chaining*) o inverso (*backward chaining*).

5. I linguaggi di comunicazione in IA

Un programma è un elenco di istruzioni che il calcolatore segue in ordine per realizzare un compito specifico. Un programma deve essere scritto in un linguaggio di programmazione, che è un insieme di simboli che possono essere elaborati dal calcolatore. I linguaggi progettati per la programmazione in IA, vengono detti "linguaggi di programmazione per l'elaborazione simbolica".

Il linguaggio d'Intelligenza Artificiale storicamente più importante è il LISP. Il LISP, è l'acronimo di "List Processor" (elaborazione di liste).

Nel 1960 John McCarthy del MIT scrive il LISP, il linguaggio di programmazione per la manipolazione di informazioni. Per rappresentare l'associazione tra i simboli viene usata una struttura chiamata lista. L'elaborazione dei simboli in una struttura a lista, è forse il concetto più importante dell'elaborazione simbolica. Una lista, nel LISP, è una sequenza di elementi, in cui ciascuno di essi può essere una particella oppure un'altra lista: una lista può essere a sua volta una successione di liste. I dati e le istruzioni del programma in LISP sono nello stesso formato per cui l'informazione relativa alle proprietà di un oggetto, cioè la conoscenza dichiarativa, può essere integrata facilmente con l'informazione relativa alle azioni che occorre eseguire, cioè alla conoscenza procedurale. Un programma in LISP ha la possibilità di modificare le proprie istruzioni o aggiungerne di nuove. È possibile anche che un programma ne strutturi uno nuovo e questo è utile, specialmente se un programma deve affrontare un nuovo compito. Questa è una possibilità importante in quei sistemi specialistici detti «sistemi esperti», in cui è necessario esplicitare come si sono ottenute certe conclusioni di ragionamento (Winston e Horn, 1984).

Tra le aree di ricerca dell'IA uno dei più ambiti è il tentativo di creare dei programmi che permettano ai computer di comprendere il linguaggio naturale.

Gli scienziati stanno cercando d'insegnare ai computer l'emulazione di un'abilità, che quasi tutti gli umani eseguono senza alcun problema. I computer però non sono ancora in grado di comprendere il linguaggio naturale, nemmeno quanto un bambino di cinque anni. Un bambino di 5 anni è in grado di combinare più parole per formare una frase che abbia un senso compiuto ed è anche in grado di comprendere frasi pronunciate da adulti. Un computer può mettere in sequenza diverse parole, seguendo un percorso logico, ma il risultato finale può risultare incomprensibile. Garth Notley del Microprocessor Consultancy Service ha ottenuto col computer molti esempi di parlato: in alcune frasi anche molto complicate, il computer riconosce i singoli spezzoni ed è guidato da un generatore di numeri casuali che collega in forme sintattiche note e corrette, senza però conoscere l'argomento. Con il computer si possono scrivere testi, bozze, copioni, discorsi, grafici, si possono creare frasi in qualsiasi lingua, ma la produzione creativa manca di un senso. Allo stato attuale risulta ancora complesso costruire sistemi computerizzati in grado di riconoscere il parlato, perché nel segnale vocale ci sono molte meno informazioni, chiaramente udibili, di quanto non si creda: quando una persona dice una sola cosa per volta, parole isolate, si comprende il 50% del discorso. Il computer che ha molta meno esperienza, non riesce ad anticipare ciò che l'interlocutore sta per dire. Il linguaggio naturale con cui normalmente si comunica è informale e può essere estremamente ambiguo, impreciso e incompleto anche se viene scritto correttamente. L'ambiguità è ciò che provoca molto spesso una errata comunicazione tra le persone ed è uno dei principali problemi nella programmazione dei computer per la comprensione del linguaggio naturale. Alcuni dei fattori che contribuiscono all'ambiguità del linguaggio naturale sono i significati multipli che spesso le parole hanno. Non è insolito che una singola parola abbia più di un significato, come nelle seguenti frasi:

Il cavo della mano

Il cavo della luce

oppure

È scattato il cane e il fucile ha sparato

È scattato il cane quando il fucile ha sparato.

La parola "cavo" nella prima frase si riferisce a una incavatura, mentre

delle parole chiave – *Keyword analysis* – che è il metodo più semplice per analizzare il contenuto di una frase. Consiste in un programma che attraverso l'analisi delle parole-chiave scandisce il testo, individuando quelle parole che è stato programmato a riconoscere. Un modo per verificare che nessun elemento di una frase vada perduto è poter condurre una sistematica analisi della sintassi della frase, mediante la tecnica di *parsing*, metodo per scomporre una frase nelle sue parti componenti, che è l'equivalente per il computer della creazione di un diagramma della frase. Il *parsing* controlla le regolarità del linguaggio naturale ed accerta che sia avvenuta l'esatta comprensione della funzione di ciascuna parola in una frase e delle sue relazioni con le altre parole. «L'analisi sintattica di una frase è solamente il primo passo verso la comprensione della frase», nota Elaine Rich, «per la comprensione deve venire prodotta un'interpretazione semantica della frase» (1983, 320). Una grammatica semantica applica la conoscenza, relativa alla classificazione di concetti in uno specifico settore, all'interpretazione di una frase, allo scopo di analizzarla relativamente al suo significato. Roger Schank intorno al 1970 sviluppò un sistema di classificazione delle situazioni, in relazione ad un numero di concetti semplici. Schank usa la dipendenza concettuale per determinare il significato, in base a una comprensione di piani ed obiettivi. Come le persone utilizzano le aspettative per comprendere informazioni incomplete ed ambigue, analogamente un programma che tenta di determinare piani e obiettivi dalle informazioni disponibili può essere in grado di ottenere, rispetto ad un programma che non tiene conto delle aspettative, una comprensione più approfondita di un passaggio del testo. Ad esempio:

“Piero ha dato il televisore a nonno Luigi”

e

“Nonno Luigi ha chiesto il televisore a Piero”.

Il compito più difficile però da affrontare nella comprensione del linguaggio naturale è la pragmatica, cioè lo studio di ciò che le persone intendono realmente. A tal fine ad un programma sono necessarie una grande quantità d'informazioni sul suo utilizzatore, in modo da poter adattare le sue risposte ai bisogni dell'individuo. Sebbene sia importante la conoscenza delle regole formali della grammatica, della sintassi, della scrittura e della semantica, i primi sistemi per la comprensione del linguaggio naturale, basati solamente su regole linguistiche formali, come i programmi per “traduzione automatica” sviluppati nei primi anni Sessanta, produssero in realtà risultati paradossali. I tentativi di

traduzione di un testo da un linguaggio ad un altro risalgono all'inizio della storia dell'informatica. Si riteneva che se il computer avesse avuto accesso alle informazioni sintattiche e lessicali di due differenti lingue, non gli sarebbe stato difficile tradurre un testo da una lingua ad un'altra. La realtà si rivelò ben altra. L'errore più conosciuto è la traduzione della frase dall'inglese al russo:

“Lo spirito è forte, ma la carne è debole”.

Quando la frase russa fu ritradotta nuovamente in inglese, l'ambiguità del linguaggio naturale mostrò i suoi risultati:

“La vodka è forte, ma la bistecca è molle”.

Le ricerche attuali sulla traduzione automatica utilizzano le tecniche di comprensione e di generazione del linguaggio naturale, per tentare di assicurarsi che il testo venga tradotto in maniera più intelligente. Nei suoi programmi Schank riassume o parafrasa ogni passaggio di testo, con una procedura che prima legge ed analizza il testo e poi con un programma di generazione del linguaggio naturale elabora un riassunto o una parafrasi del testo.

Quando si richiama mentalmente l'immagine di un determinato oggetto, si richiama allo stesso tempo un gruppo di tipici attributi di quell'oggetto. La forma di rappresentazione della conoscenza IA che tenta di emulare questo raggruppamento di attributi viene chiamata *frame* e venne inizialmente proposta da Marvin Minsky (Minsky, 1968). Ad esempio, se un venditore di automobili dicesse:

“Ho una automobile da vendere”

la parola automobile farebbe nascere aspettative nella mente: ci si aspetterebbe di trovare un oggetto con quattro ruote, di una certa forma e dimensione, per un uso definito. I *frame* si sono resi validi in molte aree dell'IA, ma soprattutto nella comprensione del linguaggio naturale. Quando la parola “automobile” entra in un sistema di linguaggio naturale basato su *frame*, tutti gli attributi disponibili corredano il sistema di ciò che il significato “automobile” comporta.

Un altro schema di rappresentazione della conoscenza particolarmente utile nell'area della comprensione del linguaggio naturale è quello chiamato *script* e proposto da Schank (1984). Il concetto di *frame* è basato sull'assunzione che esistano una serie di aspettative riguardanti gli oggetti, mentre lo *script* comporta una certa sequenza nella successione degli eventi. Schank fa l'esempio dello “script del ristorante”: se si va al ristorante ci si aspetta di trovare certi oggetti, come tavoli,

de di processori di tipo relativamente semplice collegati fra loro dinamicamente, e che attualmente sembrano la miglior soluzione per quei problemi che richiedono grandi risorse computazionali.

Oggi in commercio sono disponibili alcune macchine parallele (oltre alla *Connection Machine* ad es. il *Transputer* e lo *N-Cube*) e la costruzione di linguaggi per la programmazione parallela così come di algoritmi paralleli sta diventando uno dei temi più corposi della ricerca informatica. Si ritiene che se un robot può vedere quello che fa, in modo da riuscire ad adattarsi ai cambiamenti che avvengono nell'ambiente, allora può diventare una macchina flessibile e utile. In questa direzione stanno proseguendo le ricerche anche di Mc Carthy e H. Lavesque (uno dei maggiori esperti del mondo di problemi legati alla rappresentazione della conoscenza secondo il primato della logica) che hanno presentato i risultati dei loro contributi sulle rappresentazioni artificiali di situazioni dipendenti dal contesto in un recente seminario, nell'ottobre 1990 presso il C.N.R. nel Secondo Convegno dell'AIIA, Associazione Italiana di Intelligenza Artificiale.

7. La robotica

Tra le tecnologie di Intelligenza Artificiale, la robotica è lo studio dei robot, macchine in grado di essere programmate per realizzare compiti manuali (Gevarter, 1985).

Davanti ad un pubblico divertito, anche se la commedia che il ceco-slovacco Karel Capek stava rappresentando era tutt'altro che allegra, faceva la sua comparsa per la prima volta, nell'anno 1921, il robot, soprannome dato a figure meccaniche che, appunto, svolgevano un "robot", ovvero un "lavoro da schiavo". «I robot universali di Rossum», questo il titolo della commedia, erano automi destinati al lavoro in fabbrica, robot con sembianze umane, quelli che più scatenano la nostra fantasia. «Herbert» è uno di loro, gli viene trovato un nome abbastanza familiare ed ha un compito ben preciso da svolgere: è stato incaricato di raccogliere le lattine vuote lasciate incustodite nei locali del laboratorio di Intelligenza Artificiale del MIT e di radunarle in un apposito contenitore, facendo molta attenzione a soppesarle ogni volta per non appropriarsi indebitamente di bibite non consumate interamente. La creazione di robot intelligenti e sofisticati, flessibili nel loro funzionamento e capaci di integrare armoniosamente operazione come mobilità, funzionamento, orientamento e riconoscimento, manipolando

o evitando oggetti, rimane una sfida aperta nel laboratorio di Rodney A. Brooks, il creatore di Herbert. Il laboratorio è lo stesso nel quale Tomaso Poggio studia le tecniche che consentiranno ai robot di vedere, muovendosi con flessibilità e consapevolezza nell'ambiente in cui operano. Grazie ad un accurato sistema di sensori, Herbert riceve dati dal mondo esterno e li ripartisce da una serie di unità di calcolo per la successiva elaborazione. Mentre i suoi occhi elettronici localizzano gli ostacoli, un apparato *laser* esplora la stanza alla ricerca delle lattine e, non appena una fotocellula gli indica il momento in cui queste vengono a trovarsi alla giusta distanza, dispositivi sensibili allo sforzo gli consentono di valutarne il contenuto e di calibrare la presa. Siamo sempre più vicini agli anneroidi della trilogia di Asimov, dotati di un cervello elettronico che ne regola il comportamento.

La differenza fondamentale tra robot non intelligenti e quelli intelligenti programmabili consiste nel fatto che i primi semplicemente eseguono dei movimenti preprogrammati, mentre nei secondi le tecniche di intelligenza artificiale impiegate per la loro programmazione permettono ai robot di comprendere l'ambiente in cui si trovano e di realizzare operazioni intelligenti e appropriate, in risposta alle vari situazioni esterne. Poiché il robot generalmente riceve le informazioni sullo stato dell'ambiente esterno da una telecamera o da qualche altro dispositivo sensore, viene spesso chiamato robot a controllo di sensori. I robot intelligenti sono attualmente in grado di formulare ed eseguire piani di azione e di controllare le proprie operazioni.

8. I Sistemi Esperti

I Sistemi Esperti rappresentano l'espressione più promettente della rivoluzione in atto negli stili e nelle tecniche che dal trattamento dell'informazione portano al trattamento della conoscenza. Uno degli aspetti di importanza centrale nelle applicazioni di Sistemi Esperti è il problema di come rappresentare in forma simbolica la conoscenza del dominio e come sviluppare su di essa attività di ragionamento. Settori applicativi come la diagnostica, la prospezione geologica, l'ingegneria genetica, l'ambito economico-finanziario, caratterizzati da problemi che richiedono decisioni da parte di esperti umani, sono entrati a far parte delle possibilità applicative dei calcolatori.

Un sistema esperto è un programma per computer progettato per comportarsi come esperto in un'area particolare (area di esperienza).

Noto anche come sistema basato sulla conoscenza, un sistema esperto tipicamente comprende un *data-base* di conoscenza, consistente di fatti circa la materia e di euristiche per applicare quei fatti.

Il sistema esperto è un sistema *software* che sotto opportune limitazioni e relativamente ad un particolare "dominio di conoscenza" simula il comportamento di un esperto umano nel risolvere problemi specialistici. I Sistemi Esperti (Cammarata, 1980) non sono né il sostituto dell'uomo né la soluzione per tutti i problemi, ma rappresentano il primo passo per trasferire la conoscenza umana di singole problematiche, intesa come cumulo di esperienze, a sistemi che sono in grado di verificare la validità dei modelli prodotti e di riapplicarli. Quello che stanno facendo i ricercatori è sistematizzare la conoscenza accumulata dai grandi scienziati e riportarla all'interno dei computer in modo che tutti possano attingervi in maniera proficua.

In Europa le attività di ricerca sono essenzialmente concentrate nel CEDIAG (*Centre du Développement de l'offre IA du Group*) in Francia e nel Centro di Pregnana in Italia; in queste strutture, oltre un centinaio di specialisti di "ingegneria della conoscenza" sono impegnati nella creazione di linguaggi, strumenti di sviluppo (Shell), nello studio di progetti *ad hoc*, nell'implementazione di tecnologia di base e nel trasferimento di *know-how* (Torriani, 1990).

Nel Secondo Convegno dell'Associazione Italiana di Intelligenza Artificiale (AIIA) (C.N.R., Roma, 16-17 ottobre 1990), M. Fox della Carnegie-Mellon University, ha presentato alcuni temi fondamentali dell'Ingegneria della conoscenza e modelli computazionali che da essa derivano per la costruzione di soluzioni nei Sistemi Esperti. Secondo Fox bisogna distinguere fra una prima e una seconda generazione di sistemi AI: quelli della prima generazione sono in genere *rule-based*, cioè basati su un sistema di regole-euristiche, nei quali si usa un' *expertise umana* o la si cattura sotto forma di regole che emulano il modo in cui la gente risolve i problemi. Questi sistemi sono abbastanza semplici in quanto è sufficiente stabilire regole dopo aver acquisito conoscenza. La tecnologia della seconda generazione utilizza invece la rappresentazione semantica e tecniche di *problem-solving* che non usano regole. Si cerca cioè di emulare il modo in cui le persone risolvono i problemi. La conoscenza viene contenuta in una rappresentazione a rete semantica, mediante una tecnica di ricerca che esegue l'identificazione del problema il *trouble shooting*: non ci vuole un tecnico di IA per inserire la conoscenza nel sistema perché esistono regole per la costruzione, reti

semantiche che permettono ad un qualsiasi esperto di un particolare ramo di costruire questa conoscenza invece di ricorrere a tecnici di IA.

Un problema fondamentale nello sviluppo dei sistemi esperti è la rappresentazione e l'uso delle conoscenze che consulenti specialisti nelle diverse aree possiedono ed usano abitualmente. La conoscenza di un esperto umano è generalmente di due tipi: la conoscenza formale riferita alla conoscenza acquisita dalla letteratura, e la conoscenza euristica che deriva dalla esperienza accumulata nella pratica di una specifica attività. Caratteristica comune a tutti i sistemi esperti è l'uso appunto di tecniche euristiche per costruire le risposte e la spiegazione di come si è giunti a una certa conclusione.

Le nozioni che stanno alla base di un sistema esperto sono abbastanza semplici; la parte più difficile è la costruzione della base della conoscenza. Occorre infatti sistematizzare le regole fattuali ed euristiche che portano alla risoluzione di un problema specifico. La base della conoscenza nasce dal lavoro interdisciplinare tra l'esperto e «l'ingegnere della conoscenza», cioè l'esperto di tecniche per rappresentare ed usare la conoscenza specifica degli esperti umani. Il lavoro necessario per ottenere un prodotto finito è lungo e difficile.

La conoscenza ha una rappresentazione in forma esplicita anziché codificata in procedure compilate, per cui il Sistema Esperto è in grado, su richiesta, di spiegare il proprio processo di ragionamento e di rendere conto all'utente delle conoscenze e dei meccanismi deduttivi utilizzati. Il meccanismo deduttivo non è deterministico ma può esplorare diverse alternative nel processo di ricerca di una conclusione, per cui soppesa, a ogni passo di funzionamento, fatti e ipotesi ed effettua la scelta più appropriata nel contesto specifico del problema che gli è presentato. La competenza degli esperti consiste in una serie di regole costituenti una metodologia volta a trovare soluzioni a problemi specifici: questa conoscenza consiste nella pratica continua di un settore e comprende metodologie risolutive rivolte alla ricerca di soluzioni quantitative, che possono avere un carattere approssimativo: è proprio questo tipo di conoscenza che caratterizza il sapere dell'esperto e che è difficile da catturare con metodologie di tipo tradizionale. Caratteristica specifica di buona parte della conoscenza di un esperto è la sua natura empirica, non formalizzata.

Il modulo più critico di un Sistema Esperto è la base di conoscenza, dalla quale dipende la qualità delle prestazioni del sistema. Una base di conoscenza contiene sia conoscenza dichiarativa (fatti su oggetti, eventi e situazioni) che conoscenza procedurale (informazione circa il corso

richiede generalmente una decisione su come cooperare. L'ultimo problema di cooperazione che insorge riguarda i servizi che la possono facilitare. Si possono "modellare" il medico generico e lo specialista come singoli sistemi esperti o "agenti" nel mondo dell'Intelligenza Artificiale. Codesti agenti dovrebbero essere in grado di risolvere il problema della collaborazione senza l'intervento umano. Una decisione di cooperare basata sulla mancanza di un'adeguata conoscenza ha due dimensioni: larghezza e profondità. Le classi di ignoranza (cioè la mancanza di conoscenza) possono essere spiegate usando il concetto di dominio. La conoscenza può essere classificata per domini di applicazione, come la giurisprudenza, la medicina, l'ingegneria, ecc.

In ogni dominio vi sono modelli e teorie che possono risolvere gamme diverse di problemi; l'ampiezza della conoscenza fa riferimento a come e quanto conosca più domini un dato esperto, mentre la profondità fa riferimento alla gamma di problemi che un esperto può risolvere entro i suoi domini di esperienza. Per descrivere la conoscenza di un sistema esperto si possono usare i domini. Un sistema esperto può decidere i limiti della propria conoscenza sulla base della ampiezza, qualora il problema assegnato non appartenga al suo dominio. Risulta invece più difficile una decisione in merito alla mancanza di conoscenza in base alla profondità: nel CES questo è ancora un problema aperto.

9. Reti neurali artificiali (Artificial Neural Systems)

Apprendere è mettere un fatto insieme ad altri fatti; noi facciamo questo con algoritmi logici. Ciò comporta applicare leggi della deduzione.

Un modo nuovo di fare Intelligenza Artificiale, oltre alle tecniche di programmazione, è quello delle reti neurali, le quali hanno la capacità di apprendere la soluzione dei problemi, mediante un *learning-system*, un addestramento o interagendo con il loro ambiente, senza una esplicita programmazione applicativa.

L'origine di questi sistemi risale all'inizio degli anni Quaranta: nel 1943, infatti, i cibernetici Warren McCulloch e Walter Pitts elaborarono modelli formali dell'attività dei neuroni, simili a quelli attualmente utilizzati. Nel 1949 lo psicologo Donald Hebb introdusse il principio del rinforzo delle sinapsi sulla base dell'uso. Nel 1960 Bernard Widrow e Marcian Hoff, due ingegneri, introdussero la regola di apprendimento che porta il loro nome. Nel 1962 Frank Rosenblatt, un informatico, descrisse le proprietà di una nuova rete, una delle prime reti neurali

artificiali, il Perceptrone, che fu duramente criticato nel 1969 da Marvin Minsky e Seymour Papert (Minsky e Papert, 1969). Il loro giudizio negativo provocò una netta caduta dell'interesse per le reti neurali. In realtà Minsky e Papert non si resero conto che il perceptrone era solo il caso più semplice di una famiglia di sistemi che poteva essere in grado di risolvere anche problemi complessi.

L'interesse per le reti neurali si è riaperto solo all'inizio degli anni Ottanta, in concomitanza con il diffondersi di nuove idee sui sistemi complessi e con il palesarsi degli insuccessi riportati in alcuni settori dall'intelligenza artificiale tradizionale.

La storia delle reti neurali può essere letta come la storia del confronto tra due concezioni diverse della mente e delle macchine destinate a riprodurla: da un lato l'Intelligenza Artificiale, dall'altro le reti neurali. L'elaborazione dati di tipo tradizionale consisteva nel risolvere problemi algoritmici, ossia problemi in cui si compie sempre un passaggio per volta: un esempio analogo è una ricetta che spiega passo per passo come arrivare bene a risultati definiti. Ma gli esseri umani pensano in modo diverso. Noi esploriamo il mondo astratto e impariamo in continuazione.

Le reti neurali non vengono programmate: esse vengono "addestrate" (Camarata, 1986). Le reti neurali artificiali sono sistemi di elaborazione dell'informazione ispirati al sistema nervoso, in grado di svolgere in modo efficiente compiti che risultano invece difficilissimi o impossibili per i sistemi normalmente implementati sui computer tradizionali, come il riconoscimento di immagini o di fenomeni, o il recupero di informazioni complete a partire da frammenti, utilizzando memorie associative.

La ricerca sulle reti neurali può fornire nuovi spunti per la comprensione del funzionamento del cervello e per la riconsiderazione delle sue presunte analogie con il computer: le reti neurali artificiali sono sistemi dotati di una architettura ispirata al cervello. Esse infatti, come il cervello, sono composte di semplici elementi di elaborazione (corrispondenti ai neuroni), provvisti di molti ingressi e di una sola uscita, disposti a strati e collegati tra loro tramite molte interconnessioni convergenti e divergenti. In particolare, il fatto che sistemi strutturalmente così semplici come le reti neurali diano luogo a comportamenti notevolmente complessi concorda con quanto si sa a proposito del cervello.

La struttura del cervello è relativamente uniforme; la sua complessità risiede nell'enorme numero di neuroni (circa 10 decimi) e nella intricatissima rete delle sinapsi (circa 10 quattordicesimi). Non è mai

Artificial Neural System). Il seminario è stato tenuto da Josef Sktzypek, dell'Università di Los Angeles. Si è rilevato che da qualche anno a questa parte si assiste ad una riscoperta delle tematiche della cibernetica ed allo sviluppo di discipline che ne proseguono i discorsi rimasti interrotti. Queste discipline vanno sotto il nome di seconda cibernetica, o di scienze della complessità, e l'attuale interesse per gli ANS, anche se non ha connessioni dirette con queste discipline, è un segno dell'espandersi di una nuova mentalità alternativa all'intelligenza artificiale convenzionale. Le proprietà che rendono gli ANS particolarmente interessanti sono rappresentate dalla possibilità di risolvere problemi nei confronti dei quali i computer si sono dimostrati deboli (visione, riconoscimento del linguaggio, ecc.), tramite implementazioni di volume, contenuto e della capacità di tollerare guasti di alcuni elementi della rete, senza interrompere il funzionamento complessivo.

Le reti neurali, che si richiamano al paradigma connessionista, sono appunto interessanti perché si tratta di strutture adattive, dotate della capacità di "apprendere" la soluzione di problemi in base alla presentazione di esempi noti, con un apprendimento supervisionato o case-based o anche attraverso l'instestazione diretta con il mondo reale con un apprendimento non supervisionato. Una rete neurale può essere implementata in vari modi. Esistono almeno tre possibili alternative: l'implementazione in silicio, l'implementazione ottica e l'implementazione biologica. Quest'ultima si basa sulle proprietà di alcune sostanze biologiche (ad es. alcune particolari proteine) ed il suo valore per il momento è esclusivamente sperimentale. L'implementazione ottica rappresenta, per alcuni aspetti, la soluzione ideale: il problema della alta connettività (la convergenza di molti segnali su un unico punto o la divergenza di un unico segnale su molti punti) è risolto molto semplicemente grazie alla propagazione dei raggi di luce nello spazio. Inoltre molte delle funzioni di trasferimento necessarie corrispondono alle funzioni di elementi ottici ben noti (determinati tipi di specchi e lenti). L'implementazione ottica può contare su un nutrito gruppo di sostenitori che stanno svolgendo interessanti ricerche, ma si ritiene normalmente che l'impiego pratico delle reti neurali richiederà dispositivi ottici integrati, che si trovano ad uno stadio di sviluppo ancora piuttosto arretrato. È quindi la via dell'implementazione su silicio ad attirare l'attenzione della maggior parte delle industrie e dei centri di ricerca. Il principale vantaggio consiste nell'impiego di tecnologie già ampiamente collaudate, mentre permangono alcuni problemi. Innanzi tutto il numero di neuroni implementabili su un singolo *chip* è limitato dall'alto numero di interconnes-

sioni, che occupano un'area considerevole. L'attività di ricerca e sviluppo per l'implementazione in silicio di reti neurali coinvolge oggi varie società e università. Tra le prime, la società Synaptics fondata da Federico Faggin, uno dei padri del microprocessore. In Italia va segnalata l'attività del gruppo del Politecnico di Torino e dell'Università di Roma 2, di cui fa parte il prof. Paserò. Vengono identificate tre principali modalità d'uso delle reti neurali:

- 1) classificazione/trasformazione di *pattern*;
- 2) ottimizzazione;
- 3) previsione.

Rientrano nel primo caso i sistemi per il riconoscimento di immagini, suoni e altri *pattern*, le memorie associative. Rientrano nel secondo caso i sistemi esperti basati su reti neurali. Un tipico problema di ottimizzazione è rappresentato dal "problema del commesso viaggiatore": dato un certo numero di città, trovare il percorso più breve per visitarle tutte. Nel terzo caso sono compresi, ad esempio, i sistemi di previsione meteorologica. Anche il campo dei cosiddetti sistemi esperti istantanei sembra essere promettente: le reti neurali possono apprendere dagli esempi e quindi possono utilizzare anche conoscenze non esprimibili sotto forma di regole "se-allora". Una rete neurale, non richiede programmazione: è in grado di imparare per semplice esposizione ad una serie di esempi forniti dall'esperto e può quindi acquisire anche quelle conoscenze non chiaramente formalizzabili.

Esiste un fervore di ricerche in aree limitrofe (perché trattano sistemi complessi non lineari e/o l'evoluzione morfologica), come la teoria del caos deterministico, la teoria dei sistemi dissipativi, la teoria delle catastrofi, la geometria dei frattali di B. Mandelbrot. È probabile che una unificazione di queste teorie o una sinergia tra esse possa produrre notevoli progressi, anche per quanto riguarda le reti neuronali. Un tentativo è già stato realizzato con i sistemi di sviluppo basati su reti neuronali.

L'idea fondamentale, è stata quella di utilizzare le proprietà collettive di sistemi composti da un gran numero di elementi computazionali semplici, per svolgere compiti di notevole complessità. Il funzionamento di questi elementi computazionali semplici è spesso ispirato da modelli semplificati dei neuroni biologici, da cui il termine Neurocomputing, con il quale si fa riferimento a tale nuova tendenza nel settore dell'informatica avanzata.

disponibili permettono solo simulazioni parziali con il cervello biologico, Pribram ha un'intuizione geniale: ricerca un'altra organizzazione che, operando insieme a quella rappresentata dai calcolatori, porti ad un'elaborazione in parallelo e che individua nei sistemi ottici. In questi sistemi la luce ha scarsa somiglianza con l'energia elettrochimica usata dal cervello e dal calcolatore per cui l'analogia considerata si riferisce ai "percorsi" che compie l'energia e alle conseguenti organizzazioni delle informazioni. Il fondamento dei sistemi ottici di elaborazione dell'informazione sta nella costruzione dell'immagine: non vi è somiglianza tra l'immagine e il programma, l'analogia sta nelle trasformazioni che avvengono tra l'immagine e il sistema ottico di elaborazione dell'informazione e tra il programma e il calcolatore. Se poi si identificano i componenti fisici da cui dipendono le trasformazioni si può costruire un modello del funzionamento del cervello. Pribram riporta una serie di esperimenti a prova della sue considerazioni. I percettroni consentono di costruire le immagini mediante un processo additivo, sequenziale e gerarchico: una figura viene fornita dalle caratteristiche dominanti che la compongono, ma manca il dettaglio. I dettagli dell'immagine si possano ottenere mediante l'utilizzo della frequenza spaziale, che consente di applicare la matematica della meccanica ondulatoria come l'analisi di Fourier e le tecniche connesse. A Cambridge il gruppo di Fergus W. Campbell (1968, 1969), e a Pisa dal gruppo di Lamberto Maffei e Fiorentini (1973), (Maffei e Mecacci, 1979) dimostrano che il processo cerebrale che può compiere le trasformazioni necessarie all'analisi della frequenza spaziale è molto complesso e va ricercato nell'ambito delle modalità di trasmissione dei segnali, entro il sistema nervoso. Gli impulsi nervosi che viaggiano lungo gli assoni realizzano questa trasmissione: le interazioni più importanti avvengono a livello soprattutto delle giunzioni sinaptiche che modulano la conduzione del segnale elettrico, mediante la creazione di potenziali lenti pre e post-sinaptici. Questa "microstruttura interattiva a potenziale lento" svolgerebbe il lavoro computazionale del cervello. Ciò viene dimostrato nella visione: ogni esperienza visiva è computata da microstrutture a potenziale lento. Questo rende plausibile l'ipotesi che la distribuzione dell'informazione nel sistema visivo dipende dalle interazioni tra i potenziali. Dato che la struttura della retina avrebbe la conformazione di porzioni lamellari del cervello come la corteccia, quanto esperito nel sistema visivo può essere generalizzato a livello di funzionamento cerebrale. Nei sistemi di elaborazione dell'informazione ottica le interazioni d'onda delle frequenze spaziali danno origine a delle configurazioni

di interferenze e alla diffusione in parallelo dell'informazione. Quando le configurazioni di interferenza vengono fissate in una registrazione permanente costituiscono degli ologrammi; analogamente viene chiamata olografica la distribuzione dell'informazione a livello cerebrale. Gli ologrammi, caratteristici dei sistemi di elaborazione ottica, costituiscono dunque un utile paradigma del funzionamento cerebrale: essi sono dotati di una notevole capacità di memoria, che proviene dalla possibilità di immagazzinare le varie frequenze in piccoli spazi e di altre proprietà come il ricordo associativo, l'invarianza di traslazione, cioè di posizione e di grandezza. La differenza però sta nel fatto che il cervello permette la costruzione di immagini e di programmi, mentre i sistemi ottici costruiscono solo le immagini. La memoria del calcolatore è organizzata in modo casuale, quella del cervello funziona secondo principi olografici: in entrambe i programmi hanno accesso ai dati memorizzati. Nel calcolatore questo accesso ai dati, all'epoca di Pribram, avviene serialmente, nel cervello in gran parte simultaneamente attraverso dei percorsi che consentono ai segnali di essere trasmessi in parallelo. Questo funzionamento simultaneo produce degli stati cerebrali temporanei simili alle configurazioni olografiche. Il funzionamento del cervello però può essere meglio definito con un modello ologonico più che olografico, perché il cervello agisce globalmente in modo "lineare", ovvero partecipa sia dei processi del calcolatore, sia di quelli dell'informazione ottica. Come il calcolatore il cervello elabora l'informazione per passi successivi, ma differisce dagli attuali elaboratori per la capacità di elaborazione parallela di gran lunga maggiore. La differenza risiede inoltre nel fatto che le regole che sottostanno all'elaborazione parallela sono più simili a quelle dei processi dell'informazione ottica che a quelle degli attuali calcolatori seriali, anche l'immagazzinamento in memoria avviene in modo olografico più che in modo casuale. «Nel funzionamento del cervello struttura e distribuzione interagiscono reciprocamente: la struttura ha la forma di programmi e la distribuzione ha la forma di ologrammi» (Pribram, 1983, tr. it. 390). Poter trasporre in IA tutto questo è l'auspicio dell'Autore.

Lo *human information processing* (Cordeschi, 1984) subisce, sotto la spinta della nuova neurofisiologia computerizzata e della nuova tecnologia informatica, una decisa trasformazione: la mente viene intesa come un sistema composto da sottosistemi distinti, autonomi e gerarchici nelle loro funzioni di elaborazione. Rispondono in parte a questa descrizione i modelli elaborati da Broadbent nel 1958, da Arkinson e Shiffrin sulla memoria nel 1968, da Sternberg nel 1969. David Marr nel

maggiori sono le istruzioni che può eseguire in un dato tempo. Pertanto, i calcolatori più veloci possono essere in grado di ottenere una soluzione euristica migliore rispetto a quelli più lenti. I programmi di Intelligenza Artificiale spesso però non richiedono calcolatori ad alta velocità, ma tendono anche a utilizzare una memoria molto più grande di altri programmi. Questa ulteriore esigenza è causata dalle caratteristiche del sofisticato ambiente *software* necessario. Inoltre i normali sistemi di trasmissione dell'informazione non sono in grado di portare i volumi di dati richiesti da un computer all'altro in tempi adeguati alle odierne capacità di elaborazione.

La tecnologia si sta muovendo rapidamente e sta utilizzando la luce come mezzo trasmissivo. Da particelle di silicio si ottiene il vetro, materiale di facile lavorazione che ha la capacità di trasmettere la luce e di variarla a seconda della forma e del colore; realizzando cavi in vetro piccolissimi chiamati fibra di vetro e costituiti da due tipi di vetro, uno per la parte interna e uno per lo strato esterno, si è visto che il raggio luminoso immesso nella fibra di vetro resta imprigionato all'interno della fibra per effetto delle riflessioni dovute allo strato esterno. La fibra ottica ha permesso la realizzazione di cavi molto sottili ma molto resistenti, che la luce attraversa con grande velocità e bassa attenuazione trasportando così un gran numero di informazioni trasformate in impulsi luminosi. Il problema attuale è che tutte le informazioni luminose devono essere trasformate in segnali elettrici per poter essere elaborati con i *microchip*. Il passo successivo è la creazione dei *microchip ottici* che operano direttamente sugli impulsi luminosi senza dover ricorrere a conversioni in segnali elettrici. Nell'Università di Glasgow si sta realizzando un *microchip ottico* che isola un canale di comunicazione da un fascio di luce inviato su una fibra ottica, lo elabora e lo reinserisce nel fascio di luce. Questo settore di ricerca è detto ottica integrata e il suo scopo è di sviluppare sistemi di elaborazione delle informazioni interamente basati sulla manipolazione della luce, che operano quindi al limite teorico della velocità.

Durante l'esposizione InterOP 90, che si è svolta negli Stati Uniti nel mese di ottobre 1990, è stato dimostrato che l'utilizzo delle connessioni tra concentratore e stazione potranno essere realizzate con connessioni in rame, con fibre ottiche in vetro a basso costo, con trasmettitori a bassa potenza o con fibre ottiche in plastica. L'attenzione degli scienziati è rivolta ora al computer ottico: il vantaggio di tale computer è che opera direttamente sulla luce per cui per un lavoro che richiede un'ora di elaborazione su un computer elettronico basterebbero 4 se-

condi a un computer ottico. Il principio su cui si basa il computer ottico è la bistabilità. Anche l'uomo vede attraverso un sistema di comunicazione optoelettronico: l'occhio coglie un'immagine di luce riflessa e la elabora attraverso un complicato sistema di lenti, iride e retina. Per noi l'immagine ha un senso soltanto perché viene tradotta in un flusso di segnali elettrici e trasmessa lungo il nervo ottico, un cavo quasi perfetto. I segnali vengono trasmessi poi al cervello, il nostro computer incorporato. Qui l'informazione viene elaborata dando luogo al riconoscimento dell'immagine. La tecnologia sta arrivando al punto in cui la traduzione in impulsi elettrici non sarà più necessaria: tutto si svolgerà direttamente attraverso la luce e alla velocità della luce.

Per quanto riguarda i nuovi progetti avviati nei primi anni Ottanta vi è la relazione giapponese del 1982 sui sistemi computerizzati della quinta generazione. Essa è fondamentalmente basata su modelli dell'elaborazione dell'informazione a parallelismo massivo o diffuso.

Nel 1945, il matematico americano John Von Neumann scrive un articolo intolato *Report on the EDVAC* (Electronic Discrete Variable Computer = calcolatore variabile a elettronica discreta) in cui descrive una struttura logica di elaborazione dell'informazione che è la base dello sviluppo dei calcolatori. I calcolatori moderni spesso sono chiamati "macchine di von Neumann" perché i metodi di elaborazione dell'informazione che utilizzano discendono direttamente dalle sue teorie. Un elemento importante della macchina di Von Neumann che ha resistito nel corso del tempo è il concetto di elaborazione sequenziale. Gran parte della attività di un computer consiste nel cercare dati, abbinarli e quindi sceglierli. I computer delle vecchie generazioni erano sequenziali, cioè effettuavano una operazione per volta, i computer paralleli invece dividono il problema per affrontarlo più facilmente. I computer paralleli sono costituiti da tanti *chip*; ogni *chip* è un *transputer*. Il *transputer* possiede capacità mnemoniche e di calcolo, ma la cosa che lo rende speciale è che comunica con gli altri *transputer*: ciò significa che si possono collegare tra loro per risolvere un problema difficile; è il concetto dell'elaborazione parallela.

Prima di riuscire ad ottenere una potenza di calcolo pari a quella della mente umana è necessario aumentare ancora la velocità di esecuzione dei calcoli. A tal fine si stanno sperimentando supercomputer che utilizzano superconduttori per aumentare la velocità di trasferimento dei segnali elettrici tra i vari *chip* e *chip* all'arseniuro di gallio, più piccoli di un capello umano, per aumentare la capacità di elaborazione. Si stanno già studiando i *chip* biologici, cioè *chip* di silicio ricoperti di materia

BIBLIOGRAFIA

AIIA (1990) Associazione Italiana Intelligenza Artificiale, II Convegno, Roma, 16-17 Ottobre.

«AI Magazine» (1988), vol. 9, n. 3, Winter 1988.

A.D. BADDELEY - N.O. BERNSEN (1989), *Research directions in cognitive science: An european perspective*, vol. 1: *Cognitive psychology*, Lawrence Erlbaum Associates, London.

D.H. BALLARD (1986), *Cortical connections and parallel processing: Structure and function*, «Behavioral and Brain Sciences», 9, pp. 67-120.

B.G. BARA (1977), *L'intelligenza artificiale: analisi e riproduzione di attività mentali umane*, Franco Angeli, Milano.

B.G. BARA (1978), *Intelligenza Artificiale*, «Ricerche di Psicologia», vol. 5.

A. BARR - E.A. FEIGENBAUM (1982), *The Handbook of Intelligence*, 3 voll., William Kaufman, Los Antos, CA, 1981-82, 1, p. 3; 1, p. 109; 1, p. 147; 2, pp. 229-235.

D.G. BOBROW - P.J. HAYES (1985), *Artificial Intelligence - Where Are We?*, «Artificial Intelligence», 25 Mar zo 1985, pp. 375-415.

M.J. BRADY (1977), *Intelligenza Artificiale*, «Ricerche di Psicologia», vol. 5.

Brattle Research Corporation (1990), *Artificial Intelligence and Fifth Generation Computer Technologies*, Boston, p. 5; p. 236.

G.D.A. BROWN (1990), *Cognitive science and its relation to psychology*, in «The Psychologist: Bulletin of the British Psychological Society», 8, pp. 339-343.

B.G. BUCHANAN - E.H. SHORTLIFFE (1984), *Rule-Based Expert System*, Addison-Wesley, Reading, MA, p. 3.

S. CAMMARATA (1980), *Informatica e conoscenza e Sistemi Esperti*, Etas Libri, Milano.

S. CAMMARATA (1986), *Reti Neuronal: l'altra faccia dell'Intelligenza Artificiale*, Etas Libri, Milano.

F.W. CAMPBELL (1968), *Application of Fourier analysis to the visibility of Gratings*, «Journal of Psychology», p. 197.

F.W. CAMPBELL (1969), *The Spatial selectivity of the Visual Cells*, «Journal of Psychology», p. 203.

Y.F. CHEN - A. PRAKASH - C.V. RAMAMOORTHY (1986), *Network Event Manager*, Proc. Computer Networking Symposium.

N. CAMELLI (1983), *Psicologia cognitivista*, Il Mulino, Bologna.

R. CORDESCHI (1984), *La teoria dell'elaborazione umana dell'informazione. Aspetti critici e problemi metodologici*, Editori Riuniti, Roma.

R. CORDESCHI (1989), *Cervello, mente e calcolatori: précis storico dell'intelligenza artificiale*, in AA.VV., *La fabbrica del pensiero. Dall'arte della memoria alle neuroscienze*, Electa, Milano.

DEMPSEY (1988), *Why Designers of FMS Should Consider Integration*, «The FMS (Flexible Manufacturing Systems) Magazine», Aprile.

H. DREYFUS (1972), *What Computer Can't Do: The Limits of Artificial Intelligence*, Harper & Row, New York.

E. FEIGENBAUM (1988), *Expert systems: Present and future*, Key-note Speech, The seventh National Conference on Artificial Intelligence, Minneapolis, MN.

E.A. FEIGENBAUM - P. MCCORDUCK (1984), *The Fifth Generation*, Signet, New York.

E.A. FEIGENBAUM - J. FELDMAN (1963), *Computers and Thought*, McGraw-Hill, New York.

J.A. FODOR (1983), tr. it. *La mente modulare*, Il Mulino, Bologna 1988.

M.S. FOX (1984), *Artificial Intelligence in Manufacturing*, Intelligent Systems Laboratory Robotics Institute, Carnegie-Mellon University, Pittsburgh.

H. GARDNER (1985), *The Mind's New Science. A history of the cognitive Revolution*, Basic Books, New York.

W.B. GEVARTER (1985), *Intelligent Machines*, Prentice-Hall, Englewood Cliffs, NJ.

K.H. PRIBRAM (1976), *Languages of the Brain*, tr. it. *I linguaggi del cervello*, Franco Angeli, Milano.

K.H. PRIBRAM (1983), *Olonomia e struttura nell'organizzazione della percezione*, in N. CARAMELLI (a cura di), *La psicologia cognitivista*, Il Mulino, Bologna.

C.V. RAMAMOORTHY - S. SBKIHAR (1987), *Coop: an Environment for Building Cooperating Expert Systems*, Computer Science Division, Univ. of California, Berkeley.

D. RANDALL - B.G. BUCHANAN (1984), *Meta-Level Language*, in B.G. BUCHANAN - E.H. SHORTLIFFE *Rule-Based Expert Systems*, Addison-Wesley, Reading, MA.

E. RICH (1983), *Artificial Intelligence*, McGraw Hill, New York.

D. RIRCHIE (1984), *The Binary Brain*, Little Brown, Boston.

F. ROSE (1984), *Into the Heart of the Mind*, Harper & Row, New York.

R.C. SHANK (1984), *The cognitive computer*, Addison-Wesley, Reading, MA, trad. it. *Il computer cognitivo*, Giunti-Barbera, Firenze 1989.

C.E. SHANNON (1953), *Computers and Automata*, Proceedings of the IRE.

H.A. SIMON (1973), *Le scienze dell'artificiale*, ISEDI, Milano.

P. SMOLENSKY (1987), *Connectionist AI, symbolic AI, and the brain*, «Artificial Intelligence Review», 11, pp. 95-109.

P. SMOLENSKY (1988), *On the proper treatment of connectionism*, «Behavioral and Brain Science», 11, pp. 1-74.

V. SOMENZI - R. CORDESCHI (a cura di) (1986), *La filosofia degli automi. Origini dell'intelligenza artificiale*, Boringhieri, Torino.

M. STEFIK ET AL. (1983), *The Architecture of Expert System*, in D.A. WATERMAN - B.D. LENAT, *Building Expert System*, Hayes-Roth Frederick, Addison-Wesley, Reading, MA, p. 103.

G. TORRIANI (1990), *Cresce il sistema informativo*, Computerworld, Speciale IA, aprile.

A.M. TURING (1950), *Computing Machinery and Intelligence*, «Mind», 59.

P. VILLERS (1984), *Intelligent Robots: Moving toward Megassembly*, in P.H. WINSTON - K.A. PRENDERGAST (a cura di), *The AI Business*, MIT Press, Cambridge, MA.

P.H. WINSTON (1984), *Perspective*, in P.H. WINSTON - K.A. PRENDERGAST (a cura di), *The AI Business*, MIT Press, Cambridge, MA, p. 1.

P.H. WINSTON - P.B.H. KLAUS (1984), *LISP*, Addison-Wesley, Reading, MA, p. 17.