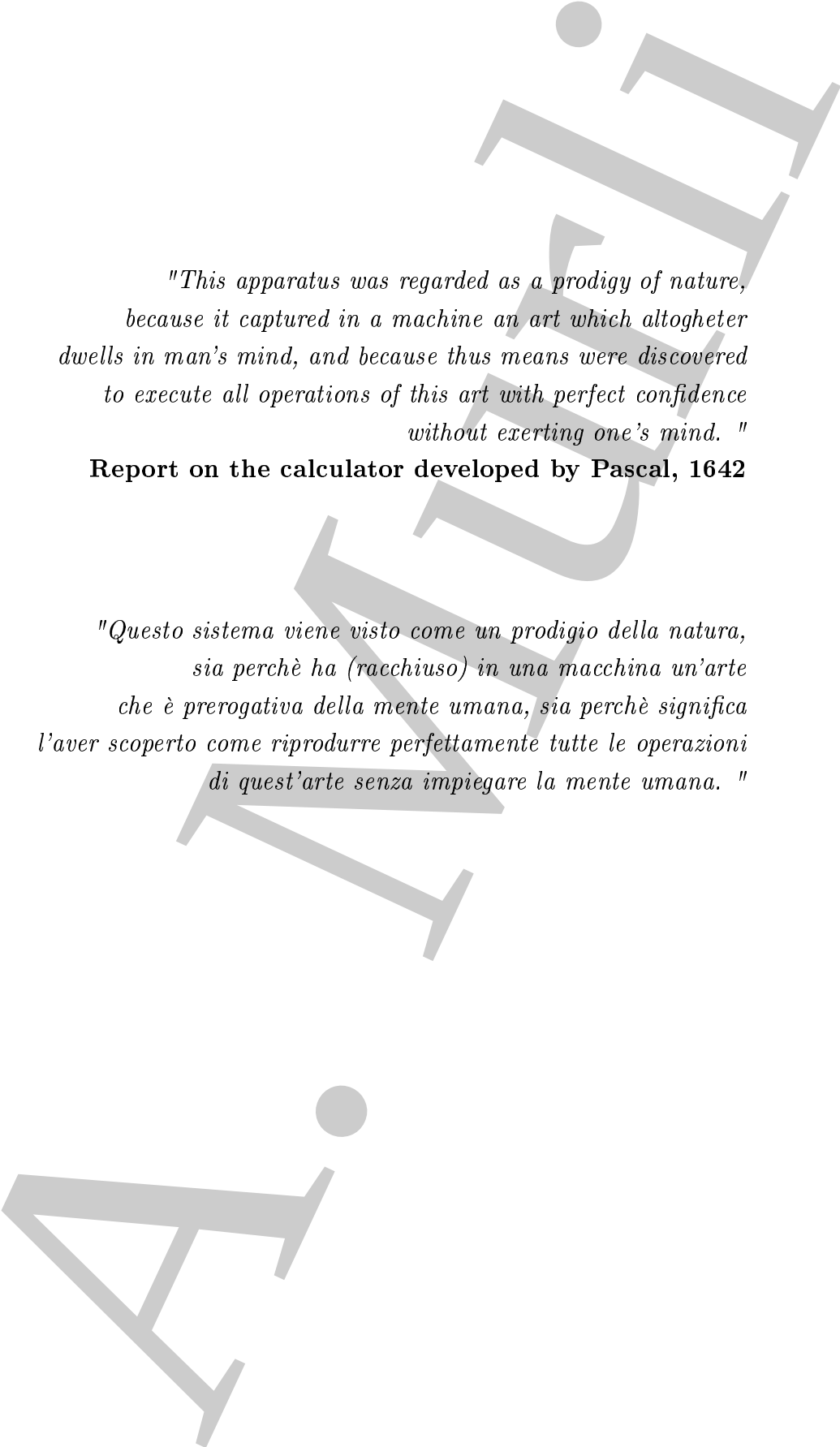




*"This apparatus was regarded as a prodigy of nature,
because it captured in a machine an art which altogether
dwells in man's mind, and because thus means were discovered
to execute all operations of this art with perfect confidence
without exerting one's mind. "*

Report on the calculator developed by Pascal, 1642

*"Questo sistema viene visto come un prodigio della natura,
sia perchè ha (racchiuso) in una macchina un'arte
che è prerogativa della mente umana, sia perchè significa
l'aver scoperto come riprodurre perfettamente tutte le operazioni
di quest'arte senza impiegare la mente umana. "*

MURLI

A.

Capitolo 6

Cenni sull'evoluzione dei calcolatori presenti sul mercato

Negli ultimi dieci anni le prestazioni dei calcolatori sono aumentate notevolmente. Nel 1977 il calcolatore più veloce era il Cray 1 ad un processore con una velocità massima di 160 Mflops. Nel 1987, l'industria giapponese era probabilmente all'avanguardia, offrendo macchine vettoriali con una velocità massima di circa 1 Gflops e una velocità ugualmente alta per il benchmark Linpack. Attualmente le macchine a più alte prestazioni raggiungono velocità dell'ordine del Teraflops. Quindi, analizzando le velocità massime, queste sono passate da 160 Mflops a 2.000 Mflops (fattore 12) tra il 1977 e il 1987 e da 2.000 a 1.000.000 Mflops (fattore 500) tra il 1987 e il 1997 (Fig. 6.1 ¹).

¹<http://www.top500.org>

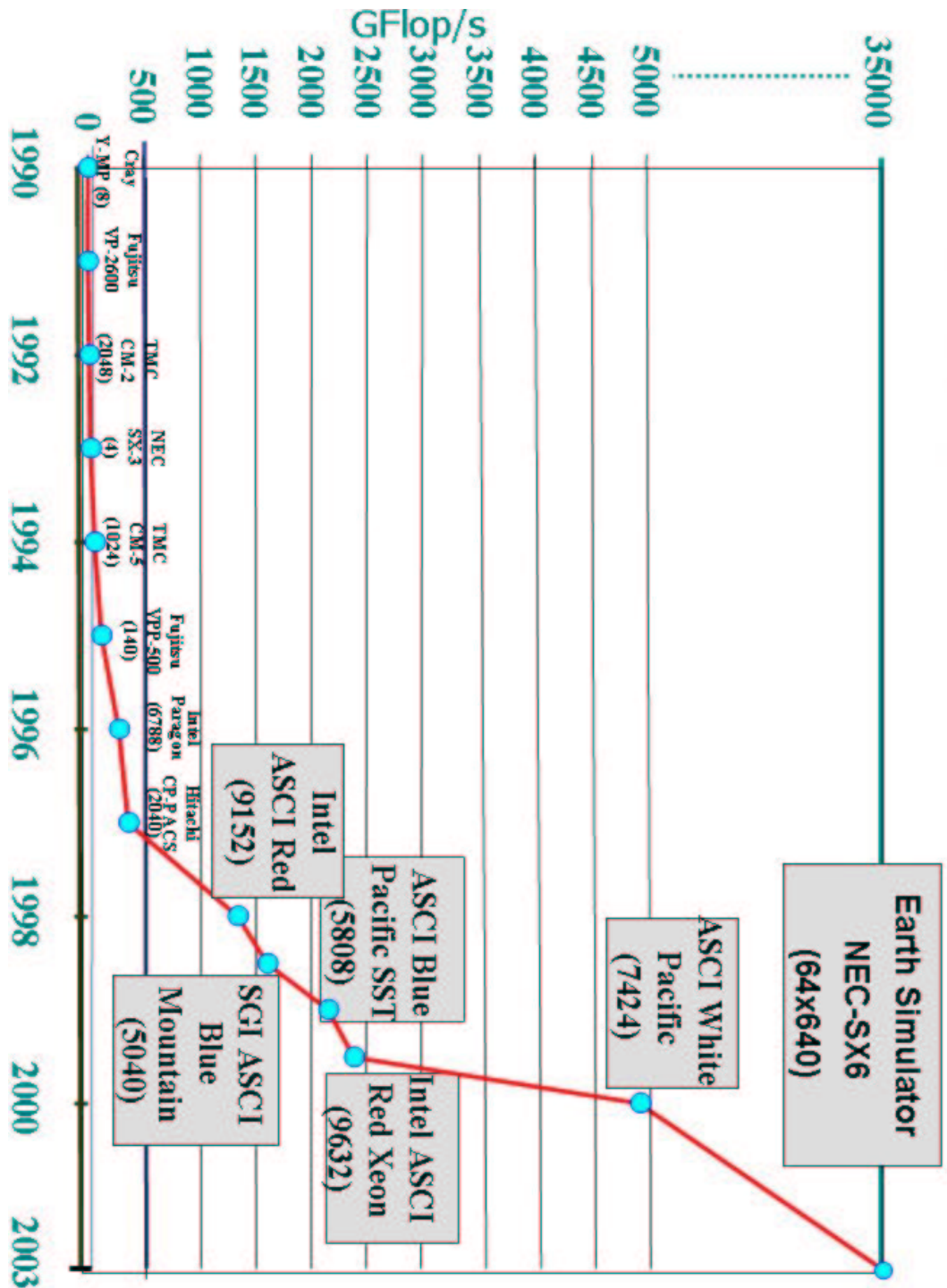


Figura 6.1: Le macchine a più alte prestazioni degli ultimi dieci anni. Un calcolo che nel 1990 veniva completato in un anno può essere oggi risolto in un'ora.

Tecnologia RISC e CISC

Le tecnologie alla base dell'architettura di un microprocessore sono la CISC (Complex Instruction Set Computer) e la RISC (Reduced Instruction Set Computer). La differenza fondamentale tra i microprocessori CISC e i microprocessori RISC risiede nel tipo di istruzioni di base (ISA): i primi possono riconoscere ed eseguire una grande quantità di istruzioni complesse, mentre i secondi riconoscono ed eseguono una quantità molto minore di istruzioni ma le eseguono più velocemente e hanno costi inferiori.

Negli anni '70 quando la memoria veniva misurata in KB ed era molto costosa, era CISC la tecnologia predominante, perché permetteva di utilizzare la memoria in modo ottimale e senza sprechi. All'interno di un'architettura CISC come quella di Intel x86, introdotta nel 1978, si possono trovare centinaia di istruzioni: semplici direttive per eseguire somme, memorizzare dati e visualizzare risultati. Nell'ambito di un'architettura di questo tipo, se tutte le istruzioni avessero la stessa lunghezza, quelle più semplici sprecherebbero memoria; le istruzioni di tipo semplice, infatti, richiedono soltanto 8 bit di memoria, mentre quelle più complesse ne richiedono 128. D'altra parte, le istruzioni di lunghezza variabile sono più complesse da elaborare per un chip, soprattutto nel caso di istruzioni CISC molto lunghe.

Negli anni '80, poi, i chip hanno raggiunto capacità sempre più elevate, mentre i loro prezzi sono crollati; è stato in questi anni che è avvenuto il sorpasso da parte della tecnologia RISC, che garantiva prestazioni decisamente migliori. Esempi di architetture RISC sono SPARC di Sun Microsystems, la serie MIPS di MIPS Technologies, Alpha di Digital Equipment, PowerPC, sviluppato da IBM e Motorola e, per finire, PA-RISC di Hewlett-Packard.

I chip RISC utilizzano istruzioni semplici e di lunghezza predefinita.

Dato che devono operare con un numero inferiore di tipologie di istruzioni, i chip RISC necessitano di meno transistor rispetto ai chip CISC e, generalmente, garantiscono prestazioni migliori a velocità di clock simili, nonostante debbano magari eseguire una quantità maggiore di istruzioni per produrre una determinata funzionalità.

La semplicità della tecnologia RISC facilita inoltre la progettazione di processori superscalari, ovvero di chip in grado di eseguire più di un'istruzione contemporaneamente (parallelismo a livello di istruzione). Attualmente quasi tutti i processori RISC e CISC sono superscalari.

Fra gli esperti la controversia sulla validità della tecnologia RISC è annosa. Alcuni affermano che i processori RISC sono più economici e più veloci. Altri fanno notare che, semplificando l'hardware, la tecnologia RISC impone un grave fardello al software. In qualche misura la discussione è oramai superata, perché i miglioramenti delle tecnologie CISC e RISC rendono i microprocessori sempre più simili. Numerosi microprocessori RISC attuali eseguono tante istruzioni quante ne eseguivano i microprocessori CISC di ieri. E nella costruzione degli attuali CISC si usano molte tecniche che prima erano associate ai soli microprocessori RISC.

6.1 Le prime macchine vettoriali e la Cray Research

Nel primo capitolo sono stati introdotti, sia pur brevemente, i calcolatori vettoriali. Tali calcolatori eseguono operazioni su vettori mediante l'utilizzo di componenti hardware opportunamente adibite a ciò.

Le unità vettoriali che eseguono le operazioni vettoriali, richiedono generalmente che gli operandi siano memorizzati in registri vettoriali. I registri vettoriali sono un insieme di registri tradizionali a cui si può accedere con una sola istruzione.

Alcuni dei vantaggi dei processori vettoriali sono:

- il numero di istruzioni è notevolmente ridotto, in quanto una singola istruzione vettoriale sostituisce un numero considerevole di operazioni aritmetiche;
- poiché le componenti di un vettore sono generalmente memorizzate in locazioni di memoria consecutive, l'accesso in memoria risulta molto più veloce;
- quando l'operando entra completamente nel registro vettoriale si riducono anche i tempi di latenza relativi a più accessi singoli.

A. Murli, *Lezioni di Calcolo Parallelo*

Versione provvisoria solo per uso personale, soggetta ad errori.

Non è autorizzata la diffusione. Tutti i diritti riservati.

Figura 6.2: *Il Cray1*

La consegna della prima macchina vettoriale, il Cray 1 (Fig. 6.2), nel 1976 a *Los Alamos Scientific Laboratory*, negli Stati Uniti, ha segnato l'inizio dell'era dei moderni "supercomputers". Il Cray 1 era una macchina caratterizzata da una nuova architettura che ha consentito l'aumento delle performance di più di un ordine di grandezza sui sistemi scalari dell'epoca. L'architettura dell'unità vettoriale del Cray 1 è stata alla base dell'intera famiglia dei sistemi vettoriali Cray fino agli anni '90. Tali macchine erano caratterizzate dall'uso di istruzioni vettoriali e di registri vettoriali. Il sistema non aveva una unità scalare separata, ma integrava efficientemente le funzioni scalari nella CPU vettoriale con il vantaggio di elevate velocità per il calcolo scalare. Il Cray 1 era non solo la più veloce macchina vettoriale, ma anche il sistema scalare più veloce dell'epoca.

Il Cray 1 fu bene accettato dalla comunità scientifica e 65 sistemi furono venduti fino al 1984, anno di fine produzione della macchina.

Alla fine degli anni settanta le maggiori industrie di calcolatori giapponesi, Fujitsu, Hitachi e NEC, iniziarono la produzione di si-

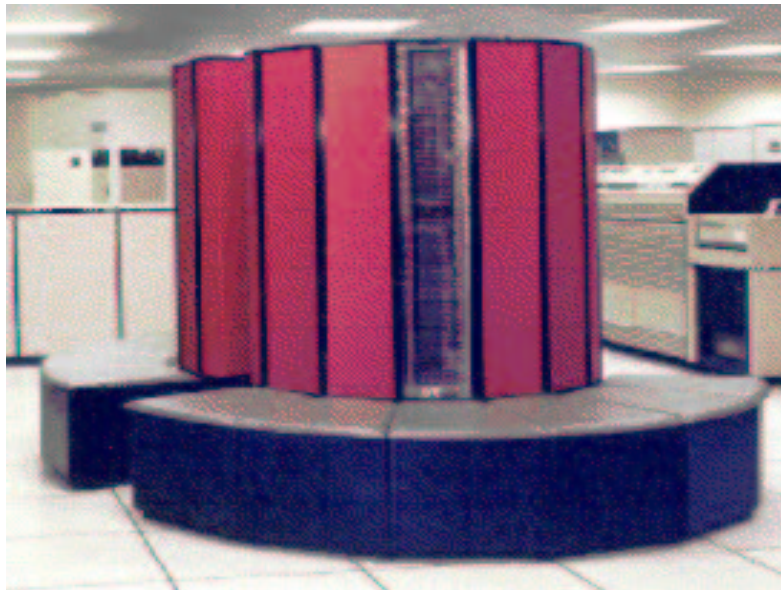


Figura 6.3: Il CrayX-MP

stemi vettoriali. La prima macchina fu distribuita sul mercato alla fine del 1983 e fu venduta principalmente in Giappone. La Fujitsu decise di espandere il proprio mercato agli USA e all'Europa tramite i suoi partners, la Amadahl e la Siemens. È per questo che alcune macchine Fujitsu, la VP100 e la VP200, erano diffuse anche in Europa. NEC provò da sola a distribuire sul mercato le proprie macchine, SX1 ed SX2, dedicandosi maggiormente alla commercializzazione. La Hitachi decise invece di non vendere all'esterno del Giappone la macchina S810. Tutte le macchine giapponesi erano caratterizzate da unità scalari e vettoriali separate e dall'uso di grandi registri vettoriali e di multiple pipelines nell'unità vettoriale.

Il passo successivo fu non solo l'aumento delle prestazioni e dell'efficienza del singolo processore, ma anche la costruzione di sistemi con più processori.

Steve Chen disegnò nel 1982 il primo sistema Cray X-MP (Fig. 6.3) con due processori. Il modello a quattro processori fu disponibili-

A. Murli, *Lezioni di Calcolo Parallelo*

Versione provvisoria solo per uso personale, soggetta ad errori.

Non è autorizzata la diffusione. Tutti i diritti riservati.

le nel 1984. I sistemi erano disegnati come sistemi multiprocessori simmetrici a memoria distribuita (symmetrical shared memory multi processor systems). Grande attenzione fu data allo sviluppo di un sistema di accesso alla memoria che fosse capace di supportare multiprocessori con ampia larghezza di banda.

Contemporaneamente Seymour Cray pose l'attenzione sullo sviluppo del sistema Cray 2. La sua attenzione era focalizzata sui chip a tecnologia avanzate. Il primo modello fu distribuito nel 1985. Con 4 processori raggiungeva una peak performance di quasi 2 GFlop/s, più del doppio dell'X-MP a 4 processori, e aveva una memoria di 2GByte. La dimensione della memoria consentiva l'uso di grossi programmi non utilizzabili sugli altri sistemi allora disponibili.

D'altra parte al Cray X-MP ha fatto seguito il Cray Y-MP. La sofisticazione del sistema di accesso in memoria, è stata una delle ragioni dell'alto grado di efficienza raggiunta da questi sistemi. Con questa linea di prodotti, successivamente comprendente il C-90 e il T-90, la Cray Research ha seguito la strada dei multiprocessori cercando sempre di ottenere la massima utilizzabilità ed efficienza dei sistemi. Il Cray Y-MP, 1988, aveva fino ad 2^3 processori, il C-90, 1991, fino a 2^4 processori e il T-90, 1995 fino a 2^5 processori.

All'inizio degli anni '80 la diffusione del sistema Cray 1 era stata fortemente aiutata dall'uso del linguaggio di programmazione standard (Fortran 77). Dopo il 1984 era disponibile per tutti i sistemi Cray un sistema operativo UNIX standard, UNICOS, cosa abbastanza innovativa per i mainframe di quel tempo. Con la disponibilità di compilatori vettoriali nella seconda metà degli anni '80 sempre più venditori di software hanno iniziato a trasportare le loro applicazioni su sistemi Cray, rendendo così possibile la distribuzione delle macchine Cray sul mercato.

Nel 1989 Seymour Cray lasciò la Cray Research per iniziare la Cray Computer e per costruire il successore del Cray 2, il Cray

3. L'idea, di nuovo, era di usare chip a tecnologia avanzata, per portare al limite le prestazioni di ogni singolo processore. Le scelte fatte purtroppo non ebbero successo. Nel 1992 fu consegnato un solo sistema e l'annunciato Cray 4 non è stato mai completato.

A causa della limitata scalabilità dei sistemi vettoriali esistenti c'era un gap tra le prestazioni dei tradizionali mainframe scalari e i sistemi vettoriali della Cray. Alcune compagnie hanno iniziato così nei primi anni '80 a sviluppare i "mini-supercomputers". L'obiettivo era di raggiungere un terzo delle prestazioni delle macchine Cray, ma solo con un decimo del costo. La compagnia che ha avuto maggior successo è stata la Convex, fondata da Steve Wallach nel 1982. Il primo sistema Convex a singolo processore, il C1, fu consegnato nel 1985. Nel 1988 è stato distribuito anche il sistema multiprocessore C2. Per il software supportato questi sistemi sono stati bene accettati negli ambienti industriali e la Convex ne ha venduto più di 500 nel mondo.

Altri esempi di calcolatori vettoriali con unità funzionali pipelined degli anni '80 sono: i Siemens Vector Processor, gli Amdahl Vector Processor, l'Alliant FX, il CDC CYBER 205.

La serie dei Siemens Vector Processor System è costituita da 5 modelli: VP30-EX, VP50-EX, VP100-EX, VP200-EX, VP400-EX. Essi sono costituiti dalle seguenti unità:

- *Vector Processing Unit (VPU)*;
- *Main Storage Unit (MSU)*;
- *Vector Storage Unit (CVSU)*;
- *Channel Processor (CHP)*;
- *Operator Console e Service Processor (SVP)*.

La Vector Processing Unit (CVPU). La VPU è la componente centrale del sistema hardware delle macchine Siemens VP. Essa può essere immaginata divisa in due parti: l'unità scalare e l'unità vettoriale.

L'unità scalare (SU) esegue le istruzioni scalari e controlla la VPU. Essa controlla tutte le istruzioni e quando rileva un'istruzione vettoriale la passa all'unità vettoriale.

L'unità vettoriale (VU) esegue le istruzioni vettoriali. La sua struttura pipeline consente l'esecuzione di molte istruzioni in un ciclo. La parte centrale della VU è costituita dalla Vector Instruction Execution Unit (VE-unit). Le istruzioni vettoriali vengono eseguite dalle seguenti unità pipeline:

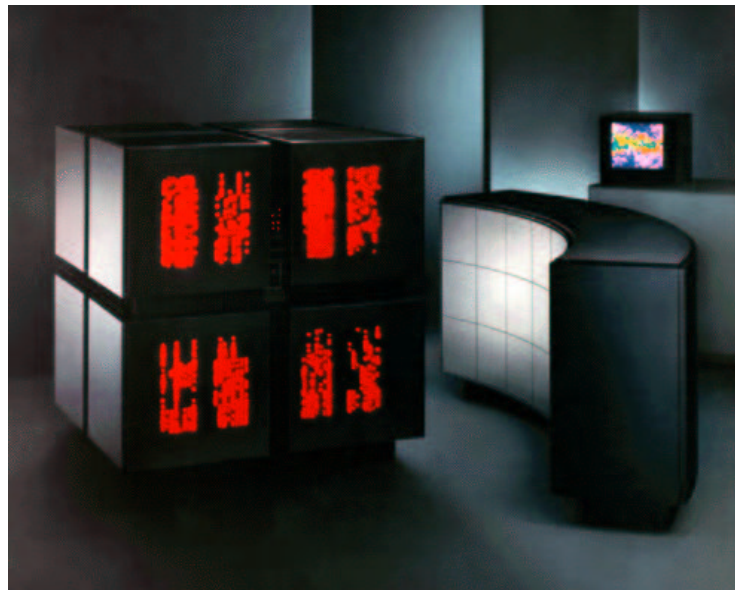
- pipeline carica/memorizza (2 per la VP100/200-EX e una per la VP30/50/400-EX);
- pipeline addizione/logica;
- pipeline multifunzionale (esegue moltiplicazioni, addizioni o moltiplicazioni per scalari e addizioni);
- pipeline per la divisione.

Le pipeline possono operare in parallelo una indipendentemente dalle altre.

La serie degli Amdhal Vectro Processor ha approssimativamente la stessa struttura di Siemens Vector Processor. In queste macchine però la Vector Unit contiene 83 istruzioni vettoriali e pipeline per le istruzioni aritmetiche.

Infine la serie delle architetture Alliant FX è costituita al più da otto elementi computazionali ognuno dotato di un processore con unità pipeline a cinque segmenti.

Un'altra categoria di calcolatori di tipo pipelined è quella dei VLIW (Very Long Instruction Word), ovvero calcolatori che usano parole di istruzione molto lunghe. In queste macchine le unità di elaborazione hanno diverse copie di ciascuna unità funzionale. Durante ogni ciclo, nel momento in cui una unità di elaborazione di tipo convenzionale leggerebbe dalla memoria una singola istruzione, le macchine VLIW sono in grado di leggerne più di una simultaneamente. Le istruzioni vengono quindi elaborate simultaneamente dalle distinte unità funzionali all'interno dell'unità di elaborazione, cosicché quest'ultima si comporta come un calcolatore parallelo.

Figura 6.4: *La Connection Machine 2*

6.2 La Thinking Machine Corporation

Nella seconda metà degli anni '80 una nuova classe di sistemi è apparsa sul mercato: calcolatori multiprocessori a memoria distribuita. Con il sostegno della *Strategic Computing Initiative del US Defense Advanced Research Agency* (DARPA, 1983) si iniziò a sviluppare questi sistemi. L'idea di base era creare sistemi paralleli senza limitazione nel numero di processori.

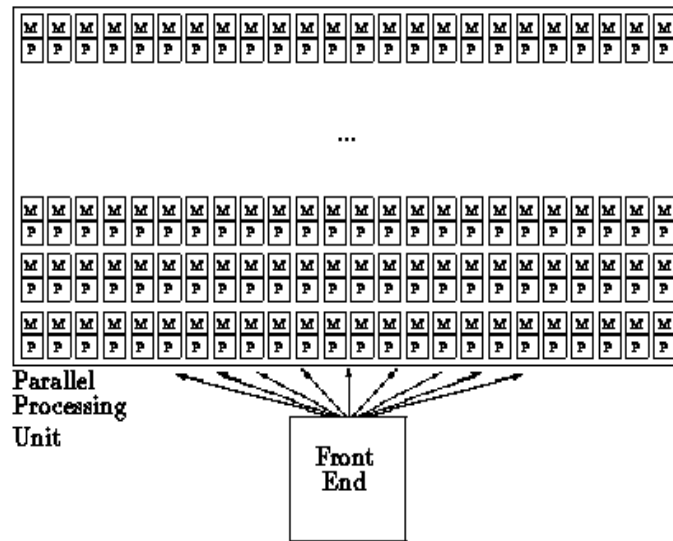
Il primo modello di sistema a parallelismo massiccio (MPP) (*Massively Parallelism Processor*) fu introdotto sul mercato nel 1985. All'inizio le architetture di differenti macchine MPP erano abbastanza diverse tra loro ad eccezione della rete di connessione, per lo più topologia ad ipercubo. La Thinking Machine Corporation (TMC) mise sul mercato la Connection Machine 1 (CM1).

Dopo la CM-1 nel 1985 la TMC introdusse la CM-2 (Fig. 6.4) che divenne il più grande sistema MPP, disegnato da Danny Hillis.

A. Murli, *Lezioni di Calcolo Parallelo*

Versione provvisoria solo per uso personale, soggetta ad errori.

Non è autorizzata la diffusione. Tutti i diritti riservati.

Figura 6.5: *La Connection Machine*

Nel 1987 la TMC iniziò a installare i sistemi CM-2. Le macchine Connection Machine erano macchine SIMD. Fino a 64K processori single bit connessi mediante una topologia ad ipercubo con $2^{11} = 2048$ unità floating point potevano lavorare insieme su di un singolo problema sotto il controllo di un singolo sistema front-end. La CM-2 fu non solo il primo sistema MPP di successo sul mercato ma anche l'unico che potesse competere con il sistema vettoriale di allora (Cray Y-MP).

La CM può essere pensata come composta da due parti principali, mostrate in Fig. 6.5.

L'unità Parallel Processing (PPU) è la parte SIMD della macchina e contiene fino a 64K processori single bit. Il Front End (FE) della macchina agisce come controllore o host per la PPU e realizza gli accessi. Il FE è generalmente una piccola workstation UNIX. Una CM può avere fino a quattro FE.

In base alla configurazione della macchina, ogni processore PPU

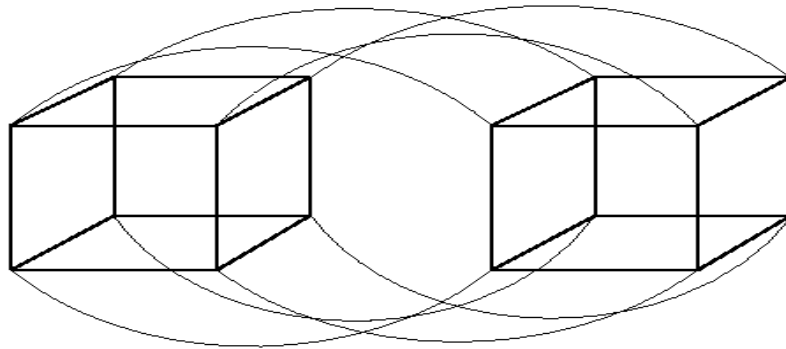


Figura 6.6: La rete di connessione della Connection Machine

ha 64K, 256K o 1024K bits di memoria, un'unità aritmetico logica (ALU), quattro registri ad un bit e un'interfaccia per due forme di comunicazione e I/O.

Oltre alla comunicazione seriale tra FE e PPU, i processori PPU possono comunicare con ogni altro in modi diversi. La forma generale di comunicazione è tramite il *router* presente su ogni chip di 2^4 processori. Questo consente ad ogni processore di comunicare con ogni altro processore ed è basata su una rete ad ipercubo come mostrato in Fig. 6.6.

L'altro modo di comunicare, più veloce, è chiamato NEWS (North-East-West-South) (Fig. 6.7) e consente ad ogni processore di comunicare velocemente con il più vicino processore. Tuttavia è limitato alla comunicazione tra processori vicini e non consente il broadcast di informazioni tra tutti i processori.

Uno svantaggio dei sistemi SIMD era la poca flessibilità dell'hardware che limitava il numero di applicazioni e programmi che potevano essere supportati. Di conseguenza la TMC decise di disegnare il nuovo sistema, la CM-5, come un sistema MIMD. Anche le altre industrie di MPP decisero di produrre macchine MIMD. La maggiore complessità di programmazione di queste macchine era più che compensata dalla maggiore flessibilità nel supportare efficientemente

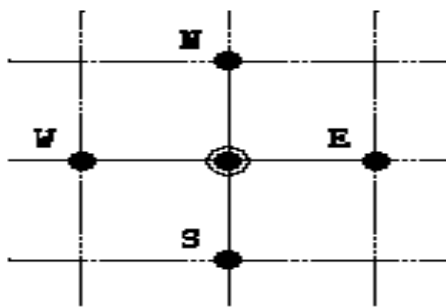


Figura 6.7: Comunicazione NEWS (North-East-West-South)

differenti paradigmi di programmazione.

6.3 Sistemi INTEL

Intel costruì i suoi sistemi basandosi su generazioni differenti dei microprocessori Intel. Il primo sistema iPSC/1 introdotto nel 1985 usava il processore intel 80286 con una rete di connessione Ethernet. Il secondo modello iPSC/2 usava il processore intel 80386 e l'iPSC/860 introdotto nel 1990 utilizzava chip i860². Per l'Intel parallelismo massiccio significava fino a 2^7 processori, limite dovuto alla dimensione massima della rete di connessione.

L'nCube, invece di utilizzare processori a basso costo, disegnò un proprio processore di tipo CISC per il sistema **nCube 10** introdotto nel 1985. Per compensare la relativamente bassa prestazione di questo processore, fu aumentato il numero massimo di processori utilizzabili. La limitazione era costituita dalla dimensione dell'iper-cubo, 12, che consentiva l'uso di massimo $2^{12} = 4096$ processori. La nuova macchina, l'**nCube 2**, introdotta nel 1989, utilizzava ancora processori propri.

²Che fu uno dei primi calcolatori paralleli in dotazione al Centro di Ricerche per il Calcolo Parallelo e i Supercalcolatori (CPS-CNR)(1990).

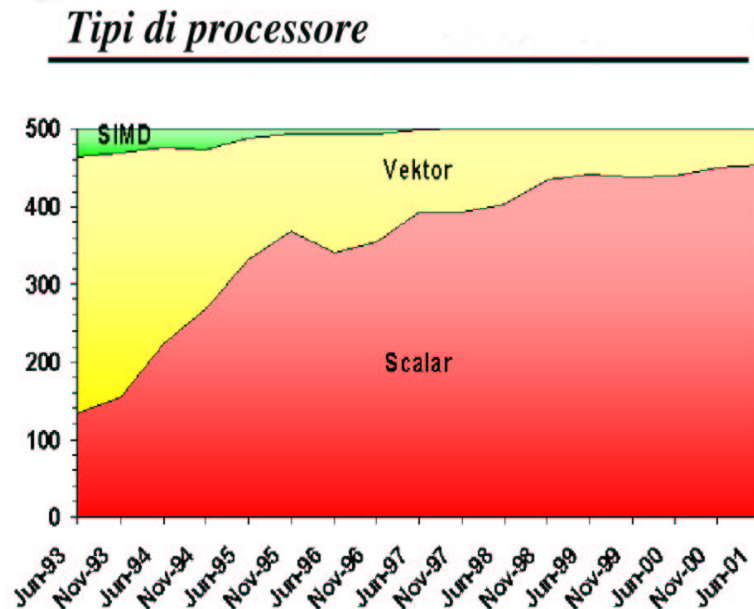


Figura 6.8: *Tipologie di CPU utilizzate come riportato dalla TOP500*

All'inizio degli anni '90 su alcune riviste comparve la frase "*The Attack of the Killer Micro*", coniata da Eugene Brook del Livermore National laboratory. Era il sinonimo della forte crescita delle prestazioni raggiungibili con i nuovi processori, dopo la rivoluzione RISC. A risentire dell'attacco furono non solo i tradizionali mainframes, ma anche i supercomputer multiprocessore vettoriali.

La nuova generazione di microprocessori relativamente economici offrivano il miglior rapporto prezzo/prestazioni. Insieme alle architetture scalabili MPP potevano essere raggiunti gli stessi livelli di prestazione delle macchine multiprocessore vettoriali ad un prezzo più basso.

Questo ha portato ad un rapido declino dei sistemi vettoriali negli anni '80, come illustrato dalla Fig. 6.8.

6.4 La TOP500

Nel 1993 Hans W. Meuer ed Erich Strohmaier ebbero l'idea della TOP500, una lista comprendente 500 calcolatori aventi le più alte prestazioni. Tutte le informazioni sulla TOP500 sono pubblicate sul sito <http://www.top500.org>.

Le statistiche sui calcolatori ad alte prestazioni sono molto utili ai produttori, agli utenti e ai potenziali utenti. In generale è utile sapere non solo il numero di sistemi installati, ma anche dove sono situati e le applicazioni per le quali vengono utilizzati. Queste statistiche facilitano le collaborazioni, lo scambio di dati e software e forniscono una migliore comprensione del mercato dei calcolatori ad alte prestazioni.

La lista viene redatta due volte all'anno da Hans W. Meuer, Erich Strohmaier, Jack Dongarra e Horst Simon. Nella lista i calcolatori vengono classificate in base alle loro prestazioni sul benchmark LINPACK.

Per la TOP500 viene usata la versione del benchmark che consente all'utente di scalare la dimensione del problema e di ottimizzare il software con lo scopo di raggiungere prestazioni più alte per un dato calcolatore.

Misurando le prestazioni per diversi valori della dimensione N del problema, l'utente può ottenere non solo la massima prestazione raggiunta R_{max} , ma anche la dimensione $N_{1/2}$ del problema che raggiunge il valore $R_{max}/2$. Questi valori, insieme alla peak performance teorica, R_{peak} , sono riportati nella lista TOP500.

Per ottenere valori uniformi delle prestazioni su calcolatori diversi, l'algoritmo usato per risolvere il sistema di equazioni nella procedura test deve essere congruo al numero standard di operazioni effettate per la fattorizzazione LU con pivoting parziale.

La Fig. 6.9 mostra i primi 10 computers della TOP500 del giugno 2001.

TOP10 21 giugno 2001

Rank	Manufacturer	Computer	Rmax	Installation Site	Country	Year	Area of Installation	# Proc	Rpeak	Nmax	N1/2
1	IBM	ASCI White, SP Power3 375 MHz	7226	Lawrence Livermore National Laboratory Livermore	USA	2000	Research Energy	8192	12288	518096	179000
2	IBM	SP Power3 375 MHz 16 way	2526	NERSC/LBNL Berkeley	USA	2001	Research	2528	3792	371712	102400
3	Intel	ASCI Red	2379	Sandia National Labs Albuquerque	USA	1999	Research	9632	3207	362880	75400
4	IBM	ASCI Blue-Pacific SST, IBM SP 604e	2144	Lawrence Livermore National Laboratory Livermore	USA	1999	Research Energy	5808	3868	431344	-
5	Hitachi	SR8000/MPP	1709.1	University of Tokyo Tokyo	Japan	2001	Academic	1152	2074	141000	16000
6	SGI	ASCI Blue Mountain	1608	Los Alamos National Laboratory Los Alamos	USA	1998	Research	6144	3072	374400	138000
7	IBM	SP Power3 375 MHz	1417	Naval Oceanographic Office (NAVOCEANO) Bay Saint Louis	USA	2000	Research Aerospace	1336	2004	374000	-
8	NEC	SX-5/128M8 3.2ns	1192	Osaka University Osaka	Japan	2001	Academic	128	1280	129536	10240
9	IBM	SP Power3 375 MHz	1179	National Centers for Environmental Prediction Camp Spring	USA	2000	Research Weather	1104	1656	-	-
10	IBM	SP Power3 375 MHz	1179	National Centers for Environmental Prediction Camp Spring	USA	2001	Research Weather	1104	1656	-	-

Figura 6.9: I primi dieci computers della TOP500 del giugno 2001

<i>Rank</i>	posizione nella TOP500
<i>Manufacturer</i>	produttore
<i>Computer</i>	tipo indicato dal produttore
<i>Installation Site</i>	utente
<i>Country</i>	paese di installazione
<i>Year</i>	anno di installazione
<i>Field of Application</i>	campo di applicazione
<i>#Proc</i>	numero di processori
R_{max}	massima prestazione
R_{peak}	peak performance teorica
N_{max}	dimensione del problema per raggiungere R_{max}
$N_{1/2}$	dimensione del problema per raggiungere $R_{max}/2$

Tabella 6.1: Nella tabella sono riportati i valori considerati per ciascun calcolatore per la formulazione della TOP500

Figura 6.10: *La Paragon XP*

Nella Tab. 6.1 sono riportati i valori considerati per ciascun calcolatore per la compilazione della TOP500.

6.5 I supercalcolatori dell'ultimo decennio presenti nella TOP500

Nel 1992 Intel inizia a produrre la serie Paragon/XP (Fig. 6.10). Basata ancora sui chips i860 la rete di interconnessione è una griglia bidimensionale che consente l'uso di $2^{11} = 2048$ processori. Nel giugno 1994 viene installato un sistema al Sandia National Laboratory che raggiungendo 143 GFlop/s sul benchmark LINPACK è la numero uno nella TOP500.

Sempre nel 1992 la TMC inizia a distribuire la CM-5, un sistema MIMD. Teoricamente può essere costruito un sistema di 16k proces-

sori, in pratica il sistema piú grande viene installato al Los Alamos Laboratory e ha 1056 processori. Raggiungeva una velocità di 60 GFlop/s. La CM-5 è stata l'ultima macchina costruita dalla TMC che nel 1994 ha terminato la produzione.

Nel giugno 1993, cinque dei primi dieci sistemi della TOP500 erano sistemi vettoriali MP. Lentamente questo numero è diminuito e l'ultimo sistema vettoriale MP a comparire nelle prime dieci macchine della TOP500 è stata nel giugno 1996 la NEC SX-4 con 2^5 processori e una velocità di 66.5 GFlop/s.

Nel 1993 IBM entra nel mercato MPP con la macchina SP1 basata sulla fortunata serie di workstation RS6000. Nel 1994 è stata realizzata l'SP2 ³ con piú nodi e maggiore velocità della rete di connessione.

Durante gli anni '90 le architetture MPP diventano sempre piú affidabili per quanto riguarda le prestazioni e l'utilizzo. In Fig. 6.11 è mostrato il cambiamento di interesse nei confronti delle varie architetture.

Il trend piú significativo è l'aumento del numero di sistemi MPP. Anche il numero di cluster aumenta e negli ultimi due anni si osserva la comparsa nella TOP500 di sistemi NOW (Network Of Workstation) basati sull'utilizzo di PC o Workstations.

6.5.1 IBM SP

L'SP (Fig. 6.12) è un sistema a memoria distribuita composto da nodi multiprocessore. Tali processori sono RISC System/6000. Sono costituiti in frames e sono interconnessi da una o piú reti di comunicazione, quali Ethernet, token ring, Opzionalmente sono interconnessi da switch ad alte prestazioni (HPS) che consentono comunicazioni veloci tra i nodi. Sono gestite da una workstation di

³In dotazione al Centro di Ricerca per il Calcolo Parallelo e i Supercalcolatori del CNR dal 1997

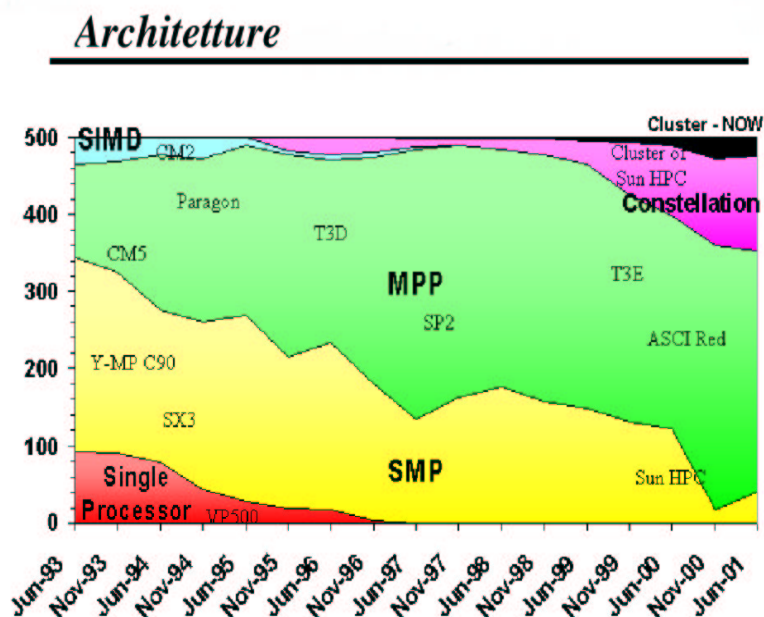


Figura 6.11: La figura mostra il cambiamento di interesse del mercato nei confronti delle diverse tipologie di architetture

controllo. Questo tipo di connessione consente l'accesso uniforme agli altri nodi per un fissato numero di nodi (aumentando il numero di nodi bisogna aggiungere altri stadi agli switches). Riduce inoltre la probabilità di congestione delle comunicazioni, in quanto realizza un *path* diretto tra ogni coppia di nodi.

Diversamente da altri produttori MPP, IBM ha posto l'attenzione sulle macchine di dimensione piccola - media.

6.5.2 CRAY T3D

Nel 1993 la Cray Research inizia finalmente ad installare macchine MPP, il Cray T3D (Fig. 6.13) basato su di una rete di connessione toroidale a tre dimensioni come specificato dal nome.

Tale macchina supporta il processore DEC chip 21064 della Di-



Figura 6.12: Una configurazione IBM SP. In basso a sinistra è rappresentata la control workstation, punto di gestione della macchina, in basso a destra sono invece rappresentate le unità di storage e di backup e in alto sono rappresentati i frames che contengono i nodi.



Figura 6.13: Il CRAY T3D

digital Equipment Corporation noto come Digital Alpha. Ogni processore ha una propria memoria locale di 16 o 64 MB. Ogni PE (Processing Element) nel Cray T3D comprende il processore DEC alpha, una memoria locale. Il PE è l'unità computazionale di base dell'architettura MIMD CRAY T3D. Un nodo CRAY T3D consiste di due PE che condividono lo switch. I PE sono connessi da una topologia toroidale 3D bidirezionale. La memoria è fisicamente distribuita tra i processori ma è globalmente indirizzabile. Ogni processore può indirizzare ogni locazione di memoria del sistema.

Per l'interfaccia di rete strettamente integrata, la latenza e la larghezza di banda della rete raggiungono valori elevati. Tuttavia le prestazioni del nodo computazionale sono fortemente limitate dalla mancanza di un secondo livello di memoria cache. La più grande configurazione installata negli USA, ha $2^{10} = 1024$ processori e ha superato la barriera dei 100 GFlop/s sul benchmark LINPACK.

6.5.3 L'HP Exemplar X-Class

Nel 1994 la Convex introduce la sua prima macchina MPP, la serie SPP1000. È una delle prime macchine a memoria virtualmente condivisa. Fino ad otto microprocessori HP sono connessi ad una memoria condivisa mediante una topologia crossbar. Diverse unità SMP sono poi connesse in un sistema a memoria distribuita. Il sistema operativo e l'hardware di connessione forniscono alla macchina una visione di memoria condivisa ad accesso non uniforme. Negli anni successivi si sono avuti vari aggiornamenti di queste macchine, l'SPP1200 nel 1995, l'SPP1600 nel 1996 e l'Spp2000 nel 1999 rinominata dalla HP Exemplar X-Class (Fig. 6.14).

L'Exemplar X-Class della Hewlett Packard fa parte della seconda generazione dei Convex SPP. È una SPP gerarchica fatta di più nodi, con ogni nodo costituito da più processori. L'Exemplar X-Class utilizza processori PA-8000 a 180 MHz con una velocità massima di

Figura 6.14: *hp Exemplar X-class*

720 Mflops. La memoria è fisicamente distribuita, ma logicamente condivisa: ogni processore o nodo può indirizzare direttamente ogni byte della memoria di ogni nodo. Ogni nodo consiste di 2^4 processori e 4 Gbytes di memoria; ogni nodo può essere considerato un SMP. Nella tassonomia di Flynn è una macchina MIMD. Al gradino più basso della scala gerarchica ci sono i singoli processori. Gruppi di processori costituiscono un nodo. Questo è il primo livello di connessione processore/memoria. I nodi sono poi connessi tra loro in una topologia ad anello o toroidale.

6.5.4 Il programma ASCI

I sistemi a più alte prestazioni degli ultimi anni e forse dei prossimi, come si evince dalla TOP500, sono certamente quelli associati all'*Accelerated Strategic Computing Initiative (ASCI)*, una collaborazione tra il Dipartimento dell'Energia americano e i tre laboratori

A. Murli, *Lezioni di Calcolo Parallelo*

Versione provvisoria solo per uso personale, soggetta ad errori.

Non è autorizzata la diffusione. Tutti i diritti riservati.

del Programma di Difesa, il *Sandia National Laboratory (SNL)*, il *Lawrence Livermore National Laboratory (LLNL)* e il *Los Alamos National Laboratory (LANL)*. Il 25 Settembre del 1995 il Presidente degli Stati Uniti, indirizzò la ricerca del Dipartimento dell'Energia nello sviluppo di attività rivolte ad assicurare un miglioramento della gestione e manutenzione sicura e fidata dell'arsenale nucleare della nazione in assenza di test nucleari. Il programma di gestione delle riserve di armi (Stockpile Stewardship Program), di cui il programma ASCI rappresenta una componente chiave, è la risposta del Dipartimento dell'Energia a questa sfida.

Il programma ASCI ha portato alla costruzione del primo computer del mondo capace di fornire prestazioni dell'ordine del teraflops; tale computer è installato al Sandia National Laboratories. Con contratti ASCI, l'IBM e la Silicon Graphics Inc./Cray Research hanno messo a punto sistemi capaci di velocità di 3 Tflops rispettivamente presso il Lawrence Livermore National Laboratory (LLNL) e Los Alamos National Laboratory. Quindi un veloce sguardo alle macchine del programma ASCI può fornire un'idea del trend a breve termine per lo sviluppo di macchine ad alte prestazioni.

ASCI WHITE

L'ASCI White (Fig. 6.15), è installata nel 2000 al Lawrence Livermore National Laboratory. È capace di performance di 10 T/Fs sul LINPACK ed è composto da $2^{13} = 8192$ microprocessori su $2^9 = 512$ nodi a memoria condivisa interconnessi mediante una connessione a grande larghezza di banda e a bassa latenza. Ogni nodo è costituito da 2^4 processori Power3-II dell'IBM. La memoria è di 6.3 terabytes e 160 terabyte di spazio distribuiti su 7000 dischi.

ASCI RED

L'ASCI RED è stata installata al Sandia National Laboratory nel



Figura 6.15: L'ASCI White



Figura 6.16: L'ASCI RED del Sandia National Laboratory

1996 (Fig. 6.16). È un sistema MIMD a memoria distribuita. Tutti gli aspetti dell'architettura sono scalabili: la larghezza di banda delle comunicazioni, la memoria principale, l'I/O. Ha 4536 nodi computazionali, ognuno con due processori Pentium Pro a 200 MHz che condividono una memoria di 128 Mbytes, per un totale di 9224 processori e 567 Gbytes di memoria. Sono previsti 76 nodi aggiuntivi destinati a supportare lavoro interattivo, I/O, reti, e funzioni del sistema operativo. I nodi sono connessi da una topologia a mesh bidimensionale.

BLUE PACIFIC

LA BLU PACIFIC è stata installata al Lawrence Livermore National Laboratory nel 1996 come seconda macchina del programma ASCI (Fig. 6.17). Consiste di due macchine separate e distinte:

- la parte aperta, accessibile per la ricerca universitaria tramite il programma ASCI;
- la parte chiusa, riservata e non accessibile dall'esterno;

La parte aperta consiste di due sistemi connessi tra loro e chiamati macchina Combined Technology Refresh (CTR). Un upgrade del 1999 ha portato la CTR ad avere una memoria di 504 GB ed un disco di 13 TB. La parte chiusa (SST) consiste di un iper-cluster di 1464 processori (SMP a quattro vie). L'iper-cluster consiste di tre settori di 488 nodi. In Tab. 6.2 sono riportate alcune delle caratteristiche della Blue Pacific.

Specifiche	SST	CTR
<i>peak performance</i>	3.9 tera OPS	892 gigaOPS
<i>memoria</i>	2.6 terabytes	504 gigabytes
<i>disco globale</i>	62.5 terabytes	10.0 terabytes
<i>disco locale</i>	17.3 terabytes	3.0 terabytes
<i>#nodi</i>	1464	320
<i>configurazione SMP</i>	4 vie	4 vie
<i>#control workstation</i>	3	1
<i>topologia</i>	omega	omega

Tabella 6.2: Nelle colonne della tabella sono riportate le Specifiche, le SST e le CTR della BLUE PACIFIC



Figura 6.17: La BLUE PACIFIC del Lawrence Livermore National Laboratory

6.5.5 L'Earth Simulator

L'Earth Simulator⁴ (ES) è un progetto delle agenzie giapponesi NASDA (National Space Development Agency), JAERI (Japan Atomic Energy Research Institute) e JAMSTEC per lo sviluppo di un'architettura a 40 *TFLOPs*. Tale progetto è nato nel Luglio 1996 e nel marzo 2000 è iniziata la sua realizzazione. Alla fine del Febbraio 2002 è iniziato il controllo sul funzionamento dell'intero Sistema Parallelo a Memoria Distribuita costituito da 640 nodi, ognuno dei quali è formato a sua volta da 2^3 processori vettoriali a memoria condivisa. L'architettura è caratterizzata da:

- 500 MHz NEC CPUs;
- 8 GFLOPs per ogni CPU, per un totale di 41 TFLOPs;
- 2Gb (4 moduli da 512 Mb FPLRAM) per CPU (in totale 10 TB);
- 640×640 switch tra i nodi;
- 16 GB/s ampiezza di banda all'interno di ciascun nodo.

Il 2 Maggio 2002 l'Earth Simulator ha ottenuto una performance di 26.58 TFLOPs mentre il suo record è stato di 35.86 TFLOPs (stime eseguite sul benchmark LINPACK). Nel Giugno 2002 ha ottenuto il primo posto nella TOP500, posto che detiene anche per il 2003.

6.5.6 NUMA-Q

Nella seconda metà degli anni '90 l'IBM ha introdotto una nuova tipologia di sistema che assicura la semplicità e la flessibilità "off the shelf" delle macchine SMP e contemporaneamente fornisce un basso tempo di latenza e alte prestazioni per applicazioni commerciali.

⁴<http://www.es.jamstec.go.jp/esc/eng/index.html>

Figura 6.18: *Earth Simulator*

Tale architettura è chiamata ccNUMA (cache coherence Non Uniform Memory Access). La prima macchina distribuita sul mercato è stata NUMA-Q nel dicembre 1996 .

L'architettura NUMA utilizza il modello SMP. NUMA realizza le sue caratteristiche di scalabilità e flessibilità connettendo più building block di SMP contenenti processori e memoria. La memoria locale su ogni building block è usata come cache per gli accessi frequenti ai dati, per limitare la frequenza degli accessi alla memoria non locale.

Il sistema IBM NUMA-Q usa l'architettura cc-NUMA. In un sistema cc-NUMA, la memoria è posizionata vicino ad un certo numero di processori, fisicamente legati insieme a formare una singola memoria. La cache coherence è riferita alla necessità di avere più copie aggiornate dei dati.

NUMA-Q è una collezione di building blocks di 4 processori Intel (quads) interconnessi mediante un insieme di collegamenti punto-

A. Murli, *Lezioni di Calcolo Parallelo*

Versione provvisoria solo per uso personale, soggetta ad errori.

Non è autorizzata la diffusione. Tutti i diritti riservati.

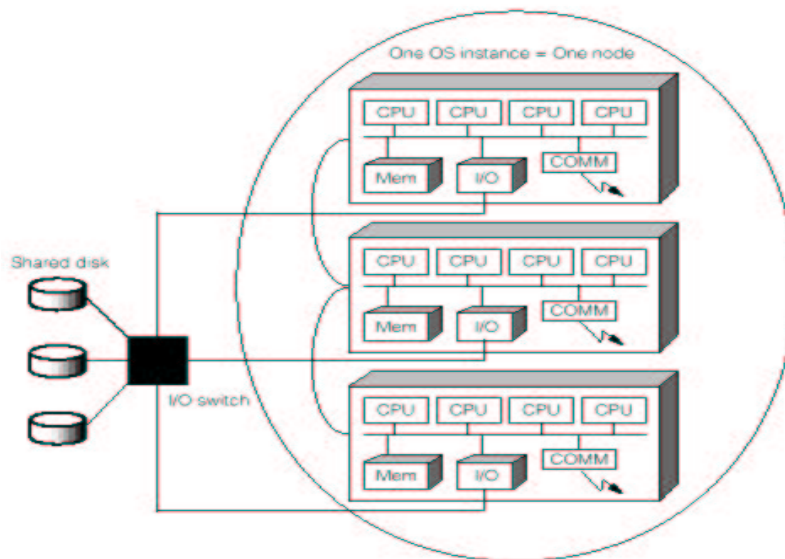


Figura 6.19: Un blocco dell'architettura IBM NUMA-Q

punto ad alta velocità. Ogni building block è una macchina SMP contenente processori, cache, memoria e bus per l'I/O. Le unità di memoria sono posizionate vicino ai processori per massimizzare le prestazioni di tutto il sistema. Anche l'I/O è vicino ad ogni processore in modo che tutti gli accessi alla memoria e all'I/O possono essere effettuati all'interno del quad.

La Fig. 6.19 mostra un blocco dell'architettura NUMA-Q.

6.5.7 SGI-ORIGIN2000

L'SGI-Origin2000 è una macchina multiprocessore ccNUMA (cache coherent Non Uniform Memory Access) disegnata e costruita dalla Silicon Graphics, INC.

Il building block di base del sistema Origin2000 è il nodo biprocessore rappresentanto in Fig. 6.20. Ogni nodo contiene fino a 4GB di memoria principale e ha una connessione diretta ad una parte

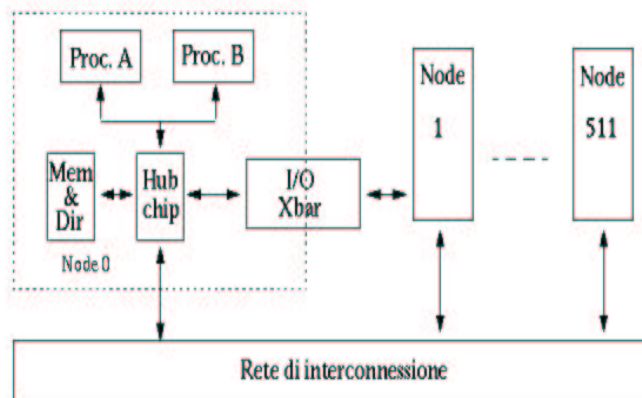


Figura 6.20: Il building block dell'Origin2000

del sistema di I/O. L'architettura supporta fino a 512 nodi con una configurazione massima di $2^{10} = 1024$ processori e 1TB di memoria principale. Sia la memoria sia il sistema di I/O sono globalmente indirizzabili.

La connessione interna al nodo è costituita da un singolo chip (Hub) che realizza un collegamento crossbar a 4 vie tra i processori, la memoria locale, l'I/O e la rete. La connessione globale è realizzata mediante rete ad ipercubo.

6.5.8 BLU MOUNTAIN

La BLU MOUNTAIN è stata installata presso il Los Alamos National Laboratory (LANL) tra giugno e novembre 1998 (Fig. 6.21). Il sistema consiste di 48 multiprocessori Silicon Graphics Origin 2000 a memoria condivisa, ognuno costituito da $2^7 = 128$ processori a 250 MHz per un totale di 6144 processori. Il cluster ha una memoria RAM di 1.5 Terabytes e ha una peak performance di 3072 Teraflops. Ogni SMP è dotato di 12 canali HIPPI bidirezionali a 800 Mbit fornendo un throughput di oltre un gigabyte al secondo.



Figura 6.21: La BLU MOUNTAIN del Los Alamos National Laboratory

6.5.9 BEOWULF

Nell'estate del 1994 Thomas Sterling e Donald Becker costruirono al CESDIS (*Center of Excellence in Space Data and Information Science*) con la sponsorizzazione del progetto ESS (*Earth and Space Science*) un cluster composto da 2^4 processori DX4. Il sistema fu chiamato *Beowulf*. Alla base della realizzazione della macchina c'era l'idea di utilizzare COTS (Commodity off the shelf), sistemi a basso costo facilmente reperibili sul mercato. La macchina ebbe un immediato successo e si diffuse velocemente negli ambienti accademici e di ricerca. L'idea di utilizzare COTS divenne così un progetto chiamato "*progetto Beowulf*"⁵ All'indirizzo URL <http://www.beowulf.org> è possibile trovare numerose informazioni e link utili sul progetto beowulf.

Il progetto Beowulf ha avuto successo per le seguenti motivazioni:

- l'industria COTS fornisce sottosistemi completamente assem-

⁵Il primo progetto Beowulf italiano è stato realizzato a Napoli nel 1996 presso il Centro di Ricerca per il Calcolo Parallelo e Supercalcolatori (CPS) del CNR.

blabili (microprocessori, schede madre, dischi, schede di rete);

- la disponibilità di opportuno software gratuito (free), in particolare il sistema operativo Linux, i compilatori GNU e le librerie di message passing MPI e PVM, che permettono di sviluppare software indipendente dall'hardware.

Una delle caratteristiche più importanti dei sistemi Beowulf è che essi realizzano il miglior rapporto costo/prestazioni.

Con la maturità e la robustezza raggiunta dal sistema operativo Linux e dal software GNU e dalla standardizzazione del message passing MPI e PVM, gli sviluppatori di software parallelo hanno una migliore garanzia di portabilità anche sui futuri sistemi Beowulf.

Nella tassonomia delle architetture parallele i cluster Beowulf appartengono ai sistemi MPP (come l'nCube, CM5, Convex SPP, Cray T3d, CrayT3E) e i NOW (Network of Workstations). Il progetto Beowulf trae vantaggi dallo sviluppo di entrambe le tipologie di architetture.

Attualmente⁶ l'installazione Beowulf con le migliori prestazioni è Locus Supercluster (Fig. 6.22) situato presso la compagnia privata Locus Discovery, Blue Bell, PA, USA. L'installazione, effettuata nel 2001, è composta da 708 nodi biprocessore per un totale di 1416 processori. 4 nodi sono Pentium III a 1000 Mhz, 1024 MBytes di memoria e un disco di 280 GBytes. Di questi 4 nodi, uno viene utilizzato per la gestione del sistema e tre per lo storage. I rimanenti 704 nodi sono utilizzati per il calcolo. Si tratta di processori Pentium III a 1000 Mhz, con 512 MBytes di memoria e un disco da 30 GBytes. Il sistema ha raggiunto una peak permonace di 1416 GFlops. La rete di interconnessione è Fast Ethernet, il sistema operativo installato è Linux e viene utilizzata principalmente per la ricerca in campo farmaceutico.

⁶Il termine "attualmente" indica il momento in cui è avvenuta la stesura del libro (Luglio 2004).



Figura 6.22: Il sistema Beowulf Locus Supercluster

MURLI

A.

Bibliografia

- [1] M. Baker, R. Buyya - *Cluster Computing at a Glance* - <http://citeseer.nj.nec.com/article/baker99cluster.html>
- [2] M.W. Berry, S.T. Dumais, G.W. O'Brien - *Using Linear Algebra for Intelligent Information Retrieval* - SIAM Review, Vol. 37, No. 4, pp.573-595, December 1995
- [3] J. Choi, J. Demmel, I. Dhillon, J. Dongarra, S. Ostrouchov, A. Petitet, K. Stanley, D. Walker, R. C. Whaley - *ScaLAPACK: A Portable Linear Algebra Library for Distributed Memory , Computers - Design Issues and Performance* - UT, CS-95-283, March 1995 - LAPACK Working Note 95, <http://www.netlib.org>
- [4] *Computer Architecture* - Computational Science Education Projects, <http://csep1.phy.ornl.gov/CSEP/CA/CA.html>.
- [5] S. Deerwester, S.T. Dumais, R.A. Harshman - *Indexing by latent semantic analysis* - Journal of the Society for Information Science, 41(6), 391-407, 1990
- [6] J. Demmel - *Lucidi del corso Applications of Parallel Computers* - Computer Science division Berkeley 1999 - http://cs.berkeley.edu/~demmel/cs267_Spr99/.
- [7] R. de Leone, A. Murli, P.M. Pardalos G. Toraldo, eds. - *High Performance Algorithms and Software in Nonlinear*

- Optimization*, Kluwer Academic Publishers, Boston ,1998, pp 1-26.
- [8] J.J. Dongarra, A.R. Hinds - *Unrolling Loops in FORTRAN* - Software - Practice and Experience, Vol. 9, 219-226 (1979)
- [9] J.J. Dongarra, J. Du Croz, S. Hammarling, R.J. Hanson - *An extended Set of Fortran Basic Linear Algebra Subprograms* - ACM Transactions on Mathematical Software, 14(1):1-17, March 1988, <http://www.netlib.org/blas/index.html>
- [10] J.J. Dongarra, J. Du Croz, I. Duff, S. Hammarling - *A set of Level 3 Basic Linear Algebra Subprograms* - ACM Transactions on Mathematical Software, 16(1):1-17, 1990, <http://www.netlib.org/blas/index.html>
- [11] J.J. Dongarra, R. van de Geijn, D.W. Walker - *A Look at Scalable Dense Linear Algebra Libraries* - UT, CS-92-155, April, 1992- LAPACK Working Note 43, <http://www.netlib.org>
- [12] J.J. Dongarra, D.W. Walker - *The Design of Linear Algebra Libraries for High Performance Computers* - UT, CS-93-188, June 1993 - LAPACK Working Note 58, <http://www.netlib.org>
- [13] J.J. Dongarra, S. Hammarling, D.W. Walker - *Key Concepts for Parallel Out-of-Core LU Factorization* - UT, CS-95-292, May 1995 - LAPACK Working Note 100, <http://www.netlib.org>
- [14] J.J. Dongarra, H.W. Meuer, H.D. Simon, E. Strohmaier - *Changing Technologies of HPC* - Future Generation Computer Systems, Volume 12, Number 5, April 1997, p 461-474.
- [15] J.J. Dongarra, H.W. Meuer, E. Strohmaier - *Top500 Supercomputer Sites* - <http://www.top500.org/> novembre 1999.

- [16] J.J. Dongarra, H.W. Meuer, H.D. Simon, E. Strohmaier - *The Marketplace of High Performance Computing* - Parallel Computing 1999.
- [17] J.J. Dongarra, V. Eijkhout - *Numerical Linear Algebra Algorithms and Software* - Journal CAM (Numerical) Linear Algebra. Volume 31(4), 28 October 1999.
- [18] J.J. Dongarra, G.E. Fagg, R. Hempel, D.W. Walker - *Message Passing Software System* - Encyclopedia of Electrical and Electronics Engineering, J. eds John Wiley & Sons, Inc., 2000.
- [19] <http://netlib.org/utk/people/JackDongarra/home.html>.
- [20] M.J. Flynn - *Some computer organisations and their effectiveness* - IEEE Trans. on Comp., C-21:948-960, 1972
- [21] L.Fosdick, E.Jessup - *An Overview of Scientific Computing.* - High Performance Scientific Computing University of Colorado at Boulder, September 1995.
- [22] I. Foster, C. Kesselman, S. Tuecke - *The anatomy of the grid* - Intl. J. Supercomputer Applications, 15(3), 2001.
- [23] G. C. Fox, P. Messina - *Architetture per i supercalcolatori.* - Le Scienze, n. 232, 1987.
- [24] E. Gallopoulos - *CSE: Content and Product* - IEEE Computational Science and Engineering, april-june 1997
- [25] A. Geist, A. Beguelin, J. Dongarra, W. Jiang, R. Manchek, V. Sunderam - *PVM: Parallel Virtual Machine. An users'-guide and tutorial for networked parallel computing* - MIT Press Scientific and Engineering Computation Janusz Kowalik, Editor, 1994.

- [26] G. Giunta, G. Laccetti, A. Murli - *Software matematico e nuove architetture* - Rivista di Informatica Vol. XXI n. 2 aprile-giugno 1991.
- [27] G. Golub, C. Van Loan - *Matrix Computations* - Second ed., Johns-Hopkins University Press, Baltimore, 1989.
- [28] J.L. Gustafson, - Re-evaluating Amadahl's Law - Communications of the ACM, 31, 532-533, 1988
- [29] J.L. Gustafson, G.R. Montry, R.E. Benner - Development of Parallel Methods for a 1024 Processor Hypercube - SIAM Journal of Scientific and Statistical Computing, 9, 609-638, 1988
- [30] Hamacher, Vranesic, Zaky - *Introduzione all'architettura dei calcolatori* - McGraw Hill 1997.
- [31] R.W. Hockney, C.R. Jessope - *Parallel Computers - Architecture, Programming and Algorithms* - Adam Hilger Ltd, Bristol, 1981.
- [32] *The IBM NUMA-Q enterprise server architecture* - http://www.sequent.com/whitepapers/numa_arch.html.
- [33] J. Laudon, D. Lenoski - *System overview of the SGI-Origin2000 Product line* - www.sgi.com.
- [34] C. Lawson, R. Hanson, D. Kincaid, F. Krogh - *Basic Linear Algebra Subprograms for Fortran usage* - ACM Trans. Math. Softw., 5:308-323, 1979
- [35] D. Marshall - Programming in C Unix System Calls and Subroutines using C - <http://www.cs.cf.ac.uk/Dave/C/CE.html>
- [36] P. Messina, A. Murli (Eds.) - *Problems and Methodologies in Mathematical Software Production* - Lecture Notes in Computer Science n. 142, Springer 1982.

- [37] P. Messina, A. Murli (Ed.) - *Practical Parallel Computing: Status and Prospectives.* - Wiley&Sons, 1991.
- [38] P. Messina, A. Murli (Ed.) - *Parallel Computing: Problems, Methods and Applications.*, Elsevier, 1992.
- [39] P. Messina - *High Performance Computers: The Next Generation (Part I)* - Computers in Physics, Vol. 11, n. 5 Sep/Oct 1997.
- [40] P. Messina - *High Performance Computers: The Next Generation (Part II)* - Computers in Physics, Vol. 11, n.6 Sep/Oct 1997.
- [41] G.E. Moore - *Cramming more components onto integrated circuits* - Electronics, Volume 38, Number 8, April 19, 1965
- [42] *MPI 1.1 - A Message Passing Interface Standard* - <http://www-unix.mcs.anl.gov/mpi/>.
- [43] A. Murli - *Lucidi delle lezioni del corso di Teoria e applicazione delle macchine calcolatrici* - Università' degli Studi di Napoli "Federico II" Dipartimento di Matematica e Applicazioni "R. Caccioppoli" 1997.
- [44] A. Murli - *Lucidi delle lezioni del corso di iCalcolo Parallelo e Distribuito, a.a. 200-2001* - Università' degli Studi di Napoli "Federico II".
- [45] A. Murli, P.D'Ambra, L.D'Amore - *Parallel Computation and problem solving methodologies: a view from some experience* -in Recent Trends in Numerical Analysis, in Advances in Computation: Theory and Practice, Edito da D. Trigiante, 2000.

- [46] M.S. Paterson, L.J. Stockmeyer - *On the number of nonscalar multiplications necessary to evaluate polynomials* - SIAM J. Comp. 2,60-66, 1973.
- [47] G.F. Carey ed. - *Parallel Supercomputing: Methods, Algorithms and Applications*, John Wiley & Sons, 1989.
- [48] *Performance of the Cray T3E Multiprocessor* - <http://www.cray.com/products/systems/crayt3e/paper1.html>.
- [49] C.J.C. Schauble - *The Connection Machine An Introduction* - HPSC Group, Department of Computer Science, University of Colorado, Boulder 1993
- [50] V. Strassen - *Gaussian elimination is not optimal* - Numerische Mathematik 13, 354-356, 1969.
- [51] A. S. Tanenbaum - *Architettura del computer: un approccio strutturale*, Ed. Jackson libri, 1996, terza edizione.
- [52] C.W. Ueberhuber - *Numerical Computation 1* - Springer, 1997
- [53] W. Ware - *The ultimate computer* - IEEE Spectrum, March 1983
- [54] R.C. Whaley. J.J. Dongarra - *Automatically Tuned Linear Algebra Software* - SC '89 Proceedings (electronic Publication), IEEE Publication, 1998, <http://netlib.org/utk/people/JackDongarra/papers.html>.