

Organização e Arquitetura de Computadores

Paralelismo em máquinas multiprocessadas

Alexandre Amory
Edson Moreno

Índice

Introdução a Classificação de Máquinas

Classificação de Flynn

SISD – Single Instruction Single Data

SIMD – Single Instruction Multiple Data

MISD – Multiple Instruction Single Data

MIMD – Multiple Instruction Multiple Data

Classificação Segundo Modelos de Memória

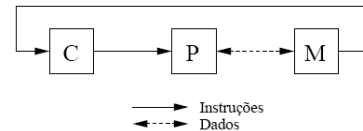
Classificação de Flynn

- Origem em meados dos anos 70
- Máquinas classificadas por **fluxo de dados** e **fluxo de instruções**

	SD (Single Data)	MD (Multiple Data)
SI (Single Instruction)	SISD Máquinas von Neumann convencionais	SIMD Máquinas Array (CM-2, MasPar)
MI (Multiple Instruction)	MISD Sem representante (até agora)	MIMD Multiprocessadores e multicomputadores (nCUBE, Intel Paragon, Cray T3D)

SISD (Single Instruction Single Data)

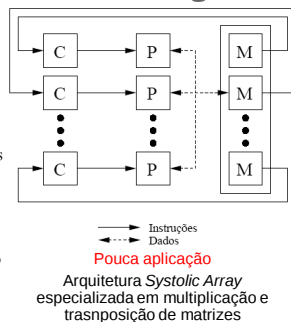
- Um único fluxo de instruções atua sobre um único fluxo de dados
- Nessa classe, são enquadradas:
 - Máquinas Von Neumann tradicionais com apenas um processador



- C – Unidade de controle
- P – Central de processamento
- M - Memória

MISD (Multiple Instruction Single Data)

- Múltiplas unidades de processamento (P), com sua unidade de controle (C), recebendo fluxo diferente de instruções
- As Ps executam suas diferentes instruções sobre mesmo fluxo de dados
- Diferentes instruções operam mesma posição de memória ao mesmo tempo, executando instruções diferentes

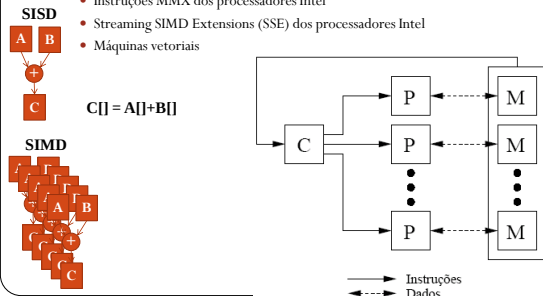


SIMD (Single Instruction Multiple Data)

- Processamento controlado por 1 unidade de controle (C), alimentada por 1 fluxo de instruções
- Mesma instrução é enviada para diversos processadores (P)
- Todos Ps executam instruções em paralelo, de forma síncrona, sobre diferentes fluxos de dados
- **Única** instrução executada sobre diferentes dados
- Computadores SIMD exploram **paralelismo em nível de dados** aplicando uma mesma operação em múltiplos itens de dados.
- SIMD pode ser muito eficiente em aplicações que possuem grande paralelismo de dados
- Exemplos de aplicações:
 - Processamento de sinais e processamento multimídia
 - Soma de vetores, multiplicação de constante com um vetor, etc

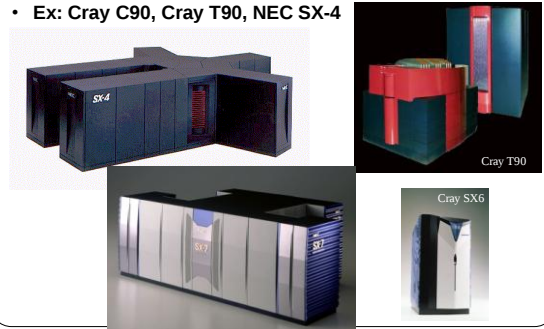
SIMD (Single Instruction Multiple Data)

- Nessa classe são enquadradas as máquinas vetoriais
 - Instruções MMX dos processadores Intel
 - Streaming SIMD Extensions (SSE) dos processadores Intel
 - Máquinas vetoriais



PVP - Parallel Vector Processor

- Custo aproximado (~1.000.000 US\$)
- Ex: Cray C90, Cray T90, NEC SX-4

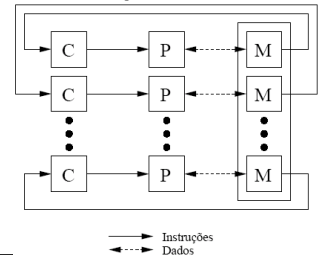


MIMD (Multiple Instruction Multiple Data)

- Cada C recebe 1 fluxo de instruções e repassa para seu P para que seja executado sobre seu fluxo de instruções
- Cada processador executa seu próprio programa sobre seus próprios dados
- Princípio MIMD é bastante genérico
 - Qualquer grupo de máquinas, se analisado como uma unidade pode ser considerado uma máquina MIMD
- Oferece mais flexibilidade
 - Um único processo usa vários processadores para aumentar o seu desempenho
 - Paralelismo em nível de thread
 - Múltiplos processos usam os vários processadores em paralelo
 - Paralelismo em nível de processos

MIMD (Multiple Instruction Multiple Data)

- Nessa classe são enquadrados
 - Servidores com múltiplos processadores (*dual, quad*)
 - Redes de computadores, cluster, computação distribuída, computação na nuvem
 - nCUBE
 - Intel Paragon
 - Cray T3D



Índice

Introdução a Classificação de Máquinas

Classificação de Flynn

Classificação Segundo Modelos de Memória

Memória Compartilhada x Não Compartilhada

Memória Centralizada x Distribuída

Distribuição de Memória

Problema: a classe MIMD é ampla demais. Ela não considera como os processadores são conectados e como a memória é vista

Uma subclassificação foi proposta, referente à localização física da memória:

- Memória distribuída (*distributed memory*)
 - Memória implementada com vários módulos
 - Cada módulo fica próximo de um processador
- Memória centralizada (*centralized memory*)
 - Memória encontra-se à mesma distância de todos os processadores → Independentemente de ter sido implementada com um ou vários módulos

Compartilhamento de Memória

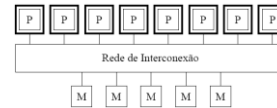
Refere-se ao espaço de endereçamento dos processadores

- Memória compartilhada (*shared memory*)
 - Único espaço de endereçamento usado para comunicação entre processadores
 - Processos acessam **variáveis compartilhadas**
 - Operações de *load* e *store*
- Memória não compartilhada
 - Múltiplos espaços de endereçamento privados → um para cada processador
 - Comunicação através de troca de mensagens → operações *send* e *receive*
- Conforme compartilhamento de memória as máquinas podem ser



Multiprocessadores

- Máquina paralela construída a partir da replicação de processadores de uma arquitetura convencional
- Todos processadores (*P*) acessam memórias compartilhadas (*M*) através de uma infra-estrutura de comunicação
 - Possui apenas um espaço de endereçamento
 - Comunicação entre processos de forma bastante eficiente (*load* e *store*)

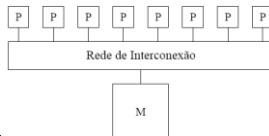


- Em relação ao tipo de acesso às memórias do sistema, multiprocessadores podem ser classificados como
 - UMA
 - NUMA, NCC-NUMA, CC-NUMA
 - COMA

Acesso Uniforme à Memória

(*Uniform Memory Access* - UMA)

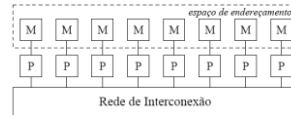
- Memória centralizada
 - Encontra-se à mesma distância de todos processadores
- Latência de acesso à memória
 - Igual para todos processadores
- Infra-estrutura de comunicação
 - Barramento é a mais usada → suporta apenas uma transação por vez
 - Outras infra-estruturas também se enquadram nesta categoria, se mantiverem uniforme o tempo de acesso à memória
- Também chamado de *symmetric (shared-memory) multiprocessors (SMPs)*



Acesso Não Uniforme à Memória

(*Non-Uniform Memory Access* - NUMA)

- Memória Distribuída
- Espaço de endereçamento único
- Memória implementada com múltiplos módulos associados a cada processador
- Comunicação processador-memórias não locais através da infra-estrutura de comunicação
- Tempo de acesso à memória local < tempo de acesso às demais → **Acesso não uniforme** → Distância das memórias variável → depende do endereço



Arquiteturas de Memória Somente com Cache

(*Cache-Only Memory Architecture* - COMA)

- Memórias locais estão estruturadas como memórias *cache*
 - São chamadas de COMA *caches*
 - COMA *caches* têm muito mais capacidade que uma *cache* tradicional
- Arquiteturas COMA têm suporte de hardware para a replicação efetiva do mesmo bloco de *cache* em múltiplos nós
 - Reduz tempo global para pegar informações

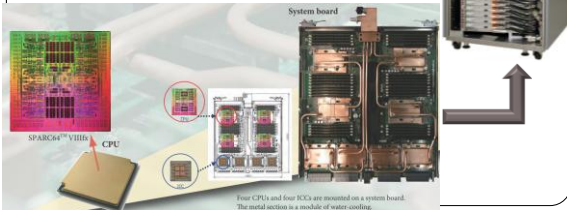


Multicomputadores

- Processadores (*P*) possuem memória local (*M*) de acesso restrito
- Memórias de outros processadores são remotas → Possuem espaços de endereçamento distintos
- Não é possível uso de variáveis compartilhadas
 - Troca de informações feita por envio de mensagens
 - Máquinas também chamadas *message passing systems*
- **Multicomputadores** são construídos a partir da replicação de toda arquitetura convencional, não apenas do processador, como nos **multiprocessadores**
- Tipicamente composto por "armários" compostos por dezenas de 'gavetas' onde cada uma possui diversos processadores

Riken, Japão

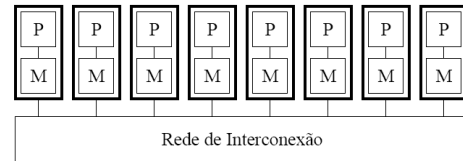
- 'Armário' do maior supercomputador do mundo (em Nov de 2011)
- Cada rack possui 24 placas com peso total de aprox 1000 kg.
- 4 CPUs e 4 ICCs por placa
- water-cooling
- CPUs SPARC64™ VIIIfx, 8 cores, 128 Gops



Sem Acesso à Memória Remota

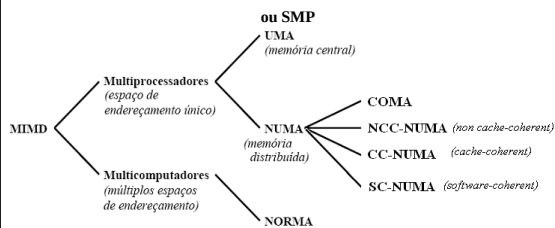
(Non-Remote Memory Access - NORMA)

- Multicomputadores são classificados como NORMA
- Também chamado de *massively parallel processors* (MPP)
- Replicação inteira da arquitetura faz que cada nó só consiga endereçar sua memória local



Visão Geral da Classificação

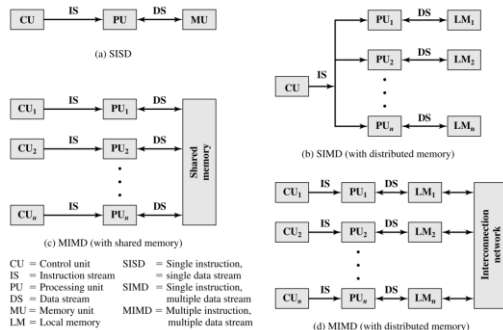
- Segundo Modelo de Memória



Duas grandes classes MIMD

- Resumo
 - MPP = muitos processadores + memória distribuída + comunicação via rede
 - SMP = poucos processadores (até poucas centenas) + memória compartilhada + comunicação via memória
 - Sistemas distribuídos = um conjunto de computadores conectados por rede que colaboram para resolver um problema

Resumão: "Big Picture"



* Computer organization and architecture: design for performance