

Arquitetura de Computadores

MiEI – 2016/17

DI-FCT/UNL

Aula 17

AC - 2016/2017

1

Desempenho dos sistemas

- Um sistema tem melhor desempenho que outro se produzir resultados em menos tempo
- A melhoria introduzida num sistema (**speedup**) é dada por:

$$S = \text{tempo antigo} / \text{novo tempo}$$

- exemplo: um programa passou a executar em 5s quando antes executava em 6s:

$$S = 6/5 = 1,2 \text{ (corresponde a 1,2x mais rápido)}$$

- $S = 1$ significa que não há alteração
- $S = 2$ significa que passou a metade do tempo

AC - 2016/2017

2

Influência dos componentes

- Componentes software:
 - Programa (algoritmos e estruturas de dados), linguagens, bibliotecas, S.O., etc...
- Componentes hardware:
 - CPU, Memória, Buses, Periféricos, etc...
- Então... exemplo:
 - Um programa não executa 2x mais rápido só porque usamos um CPU 2x mais rápido!
 - O *speedup* obtido será proporcional ao tempo de execução do CPU no tempo total do programa...
 - os tempos de espera pela memória, periféricos, etc. mantêm-se

AC - 2016/2017

3

Exemplo – o que é melhor?

- Considere um sistema que distribui o tempo total de execução nas actividades A e B:



- Se actividade B é tornada 5x mais rápida:



- Se actividade A é tornada 2x mais rápida:

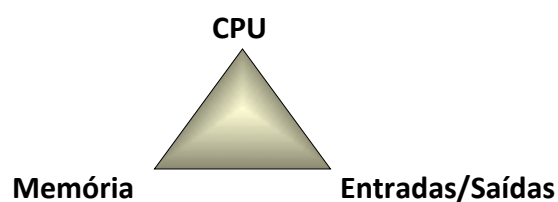


AC - 2016/2017

4

Na arquitectura de computadores

- Procura-se otimizar os vários componentes na medida do respectivo peso nos tempos de execução dos nossos sistemas

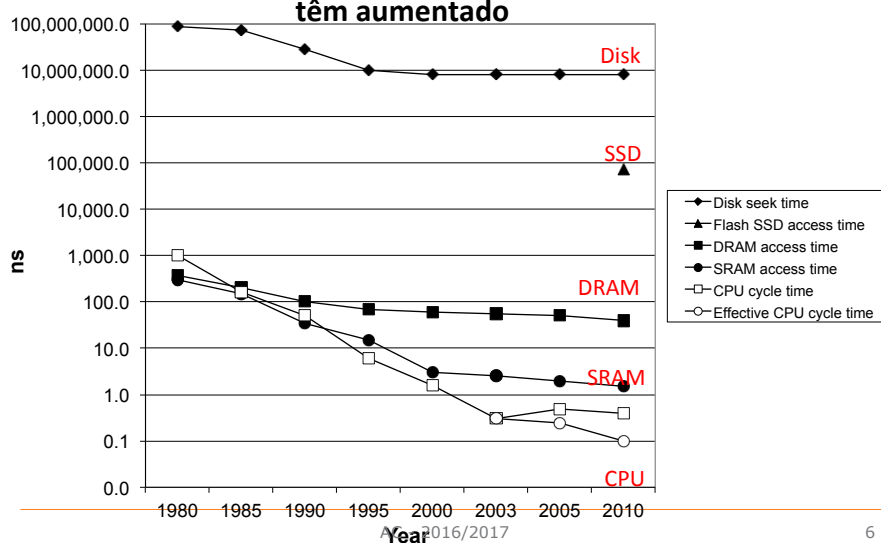


AC - 2016/2017

5

CPU-Memory Gap

As diferenças de velocidade entre RAM ou disco e o CPU têm aumentado

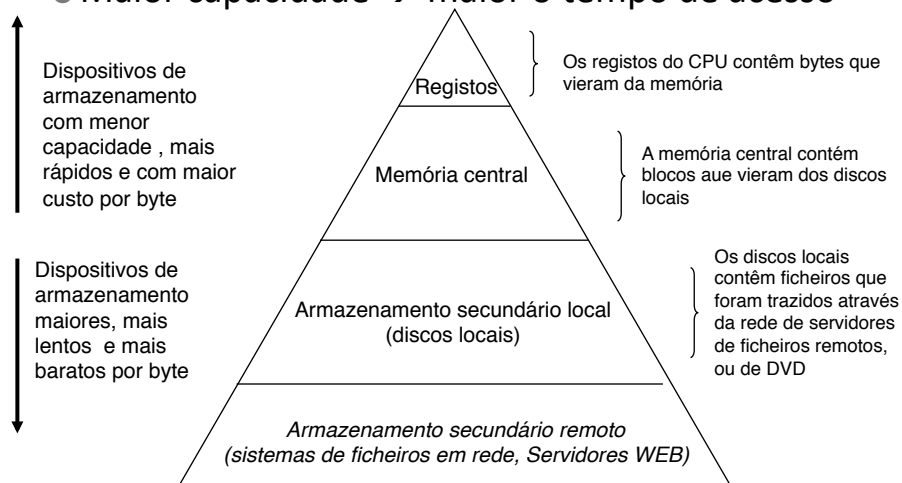


AC - 2016/2017

6

Hierarquia de níveis de memória

- Maior capacidade → maior o tempo de acesso

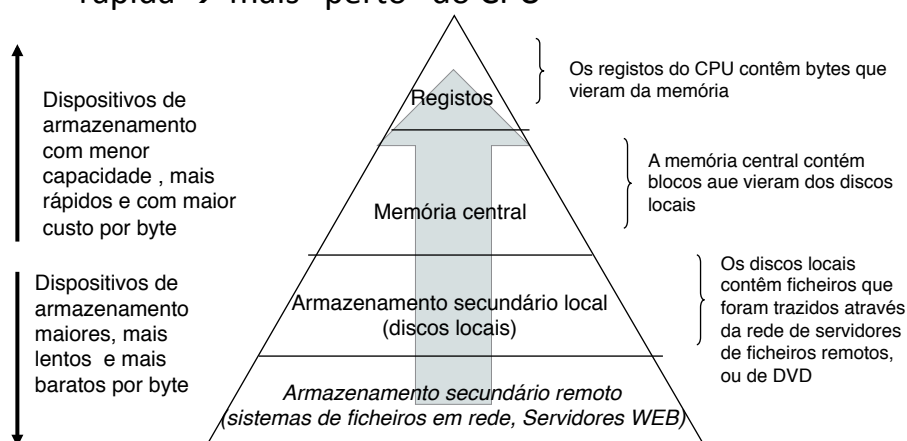


AC - 2016/2017

7

Hierarquia de níveis de memória

- Procura-se ter os dados a usar na memória mais rápida → mais “perto” do CPU



AC - 2016/2017

8

Migração dos dados para “cima”

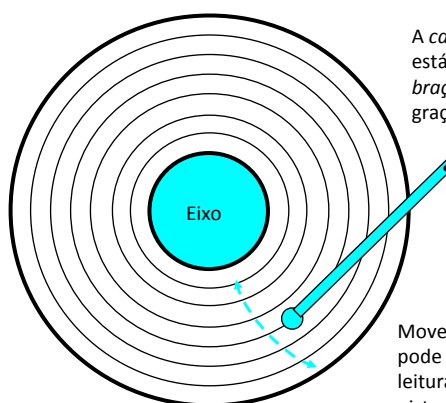
- Porquê de usar registos? Porque não executar sempre em memória?
- Porquê trazer os programas para memória? Porque não executar diretamente do disco?
- Desempenho vs. possibilidade tecnológica e custos
- O SO tenta trazer para memória o que está no disco
 - As aplicações e os dados que as aplicações leem
- O código nos programas procura tirar o melhor partido dos registos
 - Variáveis são copiadas para os registos

AC - 2016/2017

9

Operação do disco (Visão de 1 só prato)

A superfície do disco gira a velocidade constante



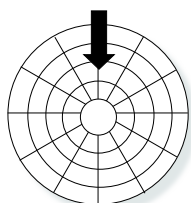
A cabeça de leitura e escrita está fixa à ponta de um braço e voa sobre a superfície graças a uma fina almofada de ar.

Movendo-se radialment, o braço pode posicionar a cabeça de leitura / escrita sobre qualquer pista

AC - 2016/2017

10

Acesso ao disco

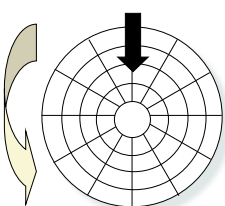


Cabeça em posição sobre uma pista

AC - 2016/2017

11

Acesso ao disco

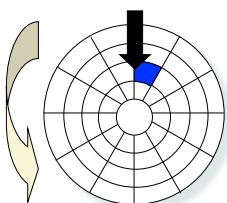


Rotação constante em sentido
contrário aos ponteiros do relógio

AC - 2016/2017

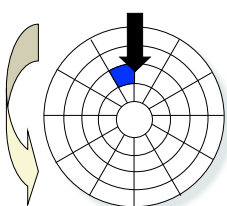
12

Acesso ao disco – Leitura



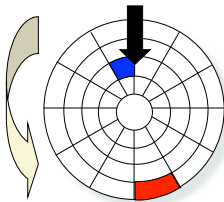
Prestes a ler o sector azul

Acesso ao disco – Leitura



Após ler o sector azul

Acesso ao disco – Leitura

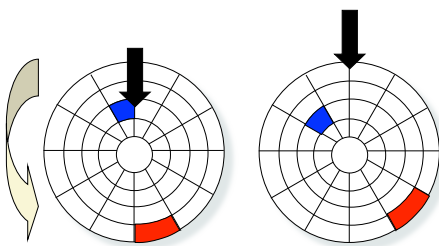


Próximo sector a ler é o sector vermelho (noutra pista)

AC - 2016/2017

15

Acesso ao disco - Seek

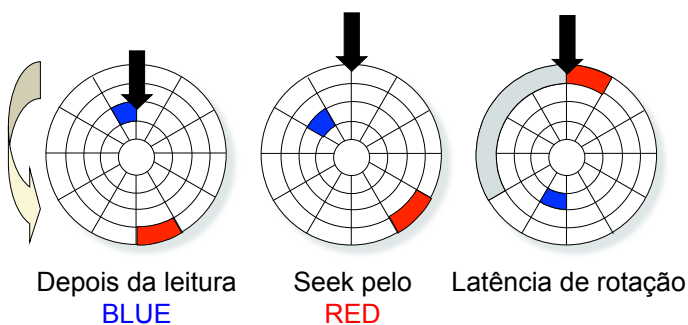


Deslocamento (Seek) para a pista do sector vermelho

AC - 2016/2017

16

Acesso ao disco – latência de rotação

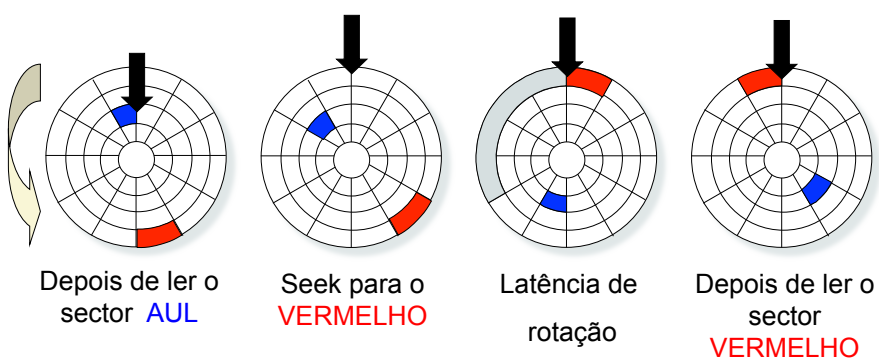


Esperar que o setor vermelho apareça debaixo da cabeça

AC - 2016/2017

17

Acesso ao disco - Leitura

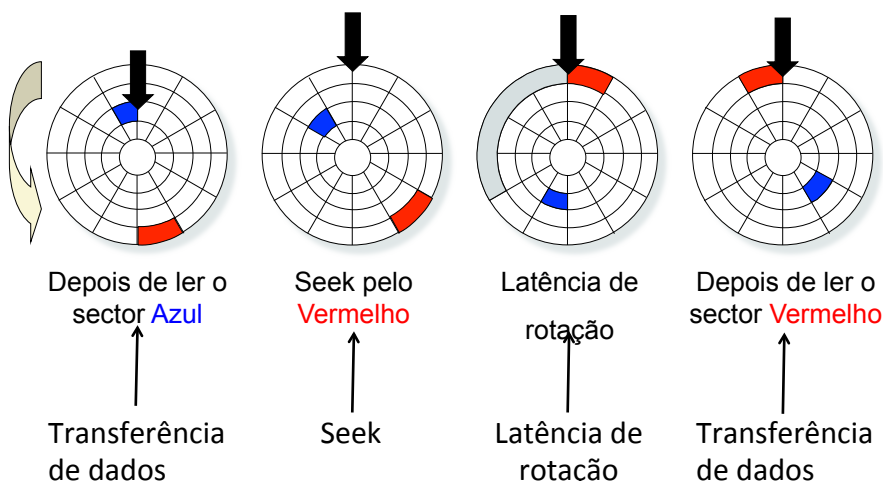


Fim da leitura do sector vermelho

AC - 2016/2017

18

Acesso ao disco – Componentes do tempo de serviço



AC - 2016/2017

19

Tempo de acesso a um setor

- Tempo médio para acesso a um dado sector pode ser aproximado por:

$$T_{\text{acesso}} = T_{\text{médio seek}} + T_{\text{médio de rotação}} + T_{\text{médio de transferência}}$$

- **Tempo de seek** ($T_{\text{médio seek}}$)

- Tempo para posicionar a cabeça sobre a pista pretendida.
- Típico: 3 - 9 ms

- **Latência de rotação** ($T_{\text{médio rotação}}$)

- Tempo de espera até que o 1º bit do sector pretendido passe debaixo da cabeça de r/w.
- $T_{\text{médio de rotação}} = \frac{1}{2} \times \frac{60}{\text{RPMs}}$
- Típica velocidade de rotação = 7200 RPMs $\rightarrow$ 4ms

- **Tempo de transferência** ($T_{\text{médio transferência}}$)

- Tempo para ler os bits do sector pretendido.
- $T_{\text{médio de transferência}} = \frac{1}{(\text{média \# sectors/pista})} \times \frac{60}{\text{RPM}}$

AC - 2016/2017

20

Exemplo de cálculo de tempo de acesso

- Dados:
 - Velocidade de rotação = 7200 RPM
 - Tempo médio de seek = 9 ms.
 - No. médio de setores/pista = 400.
- Determina-se:
 - Tmédio de rotação = $1/2 \times (60/7200 \text{ RPM}) = 4 \text{ ms}$.
 - Tmédio de transferência = $1/400 \text{ sets/pista} \times 60/7200 \text{ RPM} = 0,02 \text{ ms}$
 - **Tacesso = 9 ms + 4 ms + 0,02 ms**
- Pontos importantes:
 - Tempo de acesso dominado pelo seek time e pela latência de rotação.
 - 1º bit de um sector é o que demora mais, o resto é de "graça".
 - Acesso a sectores consecutivos poupam no seek e rotação
 - O disco pode ser 100 000x mais lento que a memória central

AC - 2016/2017

21

Blocos lógicos no disco

- Os discos modernos apresentam uma versão abstracta e simplificada da geometria do disco:
 - Os sectores disponíveis são vistos como uma sequência de blocos lógicos (0, 1, 2, ...)
 - **LBA - Logical block addressing**
- Mapeamento entre blocos lógicos e blocos físicos
 - Mantido por *hardware/firmware* no controlador do disco
 - Converte pedidos de blocos lógicos em triplos ordenados (superfície, pista, sector).

AC - 2016/2017

22

Tendências

SRAM

Métrica	1980	1985	1990	1995	2000	2005	2010	2010:1980
\$/MB	19 200	2 900	320	256	100	75	60	320x
access (ns)	300	150	35	15	3	2	1.5	200x

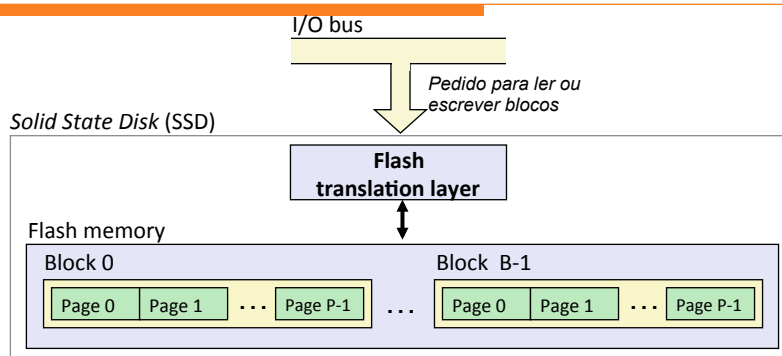
DRAM

Metric	1980	1985	1990	1995	2000	2005	2010	2010:1980
\$/MB	8 000	880	100	30	1	0.1	0.06	130 000x
access (ns)	375	200	100	70	60	50	40	9x
typical size (MB)	0.064	0.256	4	16	64	2 000	8 000	125 000x

Disk

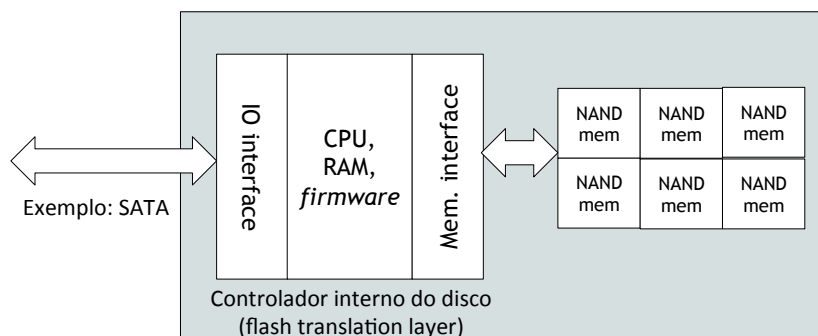
Metric	1980	1985	1990	1995	2000	2005	2010	2010:1980
\$/MB	500	100	8	0.30	0.01	0.005	0.0003	1 600 000x
access (ms)	87	75	28	10	8	4	3	29x
typical size (MB)	1	10	160	1 000	20 000	160 000	1 500 000	1 500 000x

Solid State Disks (SSDs)



- Os dados estão em páginas (eq. setores), dentro de blocos
- As escritas só são possíveis em páginas “apagadas”
- Uma página só pode ser apagada, apagando todo o bloco
- Um bloco fica inutilizado após ser apagado ~100 000 vezes.

Componentes internos ao Disco



AC - 2016/2017

25

Solid State Disks (SSDs)

Leituras sequenciais	300 MB/s	Escritas sequenciais	250 MB/s
Leituras aleatórias	170 MB/s	Escritas aleatórias	70 MB/s
Tempo de acesso leitura	20µs	Tempo de acesso escrita	200µs

- A escrita de uma página **pode** exigir apagar o bloco
- Apagar um bloco é lento (cerca de 1 ms)
- A escrita de uma página pode desencadear a cópia de todas as páginas em uso no bloco:
 1. Escolher um bloco livre e apagá-lo se necessário
 2. Copiar as páginas a manter do bloco antigo para o novo
 3. Escrever a página nesse bloco
 4. Marcar o bloco antigo como livre
- Esta gestão é feita pelo *firmware* interno ao disco

AC - 2016/2017

26

Comparação entre SSD e discos magnéticos

● Vantagens

- Sem partes móveis → maior rapidez e robustez, menos consumo

● Desvantagens

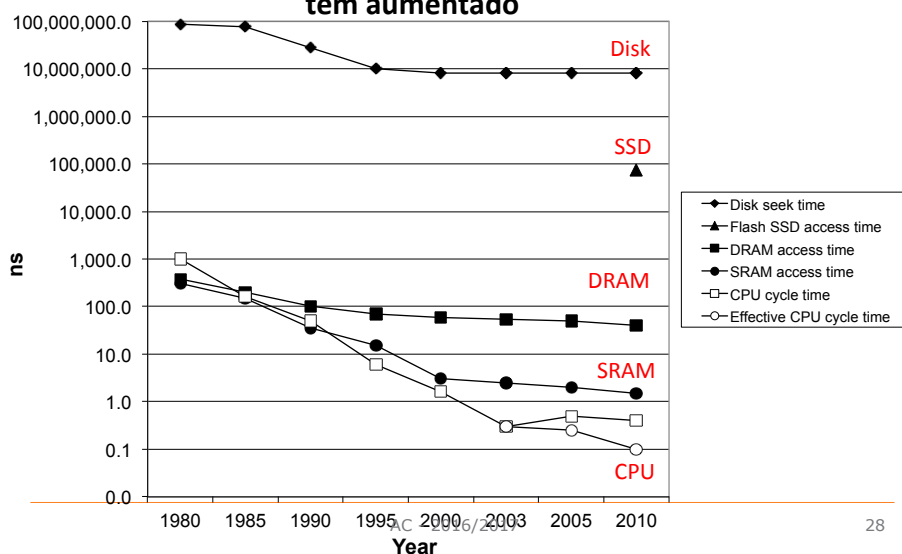
- Tem o potencial para se gastarem
 - Minimizado pela “*wear leveling logic*” na *flash translation layer*
 - E.g. Intel X25 garante 1 petabyte (10¹⁵ bytes) de escritas aleatórias até que o disco se torne inútil
- 100 x mais caro por byte (tem vindo a diminuir)

AC - 2016/2017

27

CPU-Memory Gap

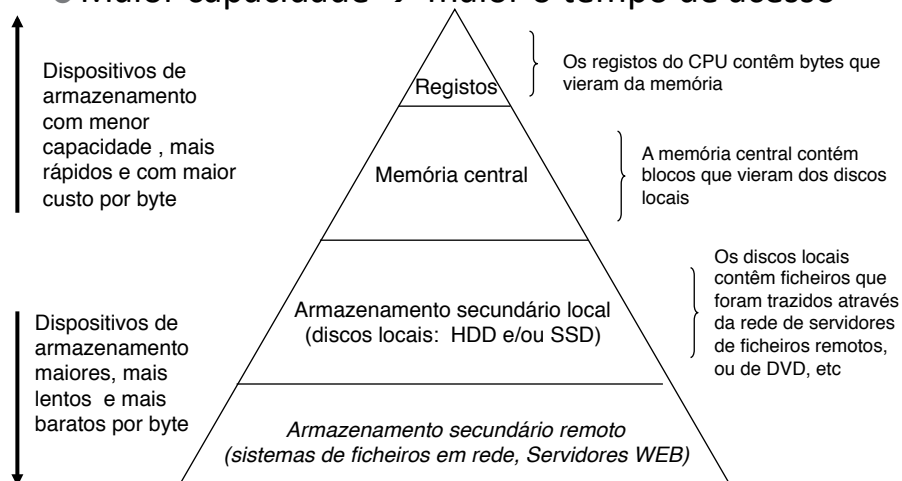
As diferenças de velocidade entre RAM ou disco e o CPU têm aumentado



28

Hierarquia de níveis de memória

- Maior capacidade → maior o tempo de acesso



AC - 2016/2017

29

Localidade

- Princípio da localidade: Os programas têm padrões conhecidos de acesso a memória
 - **Localidade temporal:** itens referenciados recentemente tendem a voltar a sê-lo.
 - **Localidade espacial:** itens com endereços próximos tendem a ser referenciados em conjunto.

Exemplo:

- **Dados**

- Elementos do array são referenciados em sequência: **Localidade espacial**
- `sum` referenciado em cada iteração:

- **Instruções**

- Instruções referenciadas em sequência: **Localidade espacial**
- Ciclo: **Localidade temporal**

```
sum = 0;
for (i = 0; i < n; i++)
    sum += a[i];
return sum;
```

AC - 2016/2017

30

Localidade (1)

- Esta função tem boa localidade?

```
int sumarrayrows(int a[M][N])
{
    int i, j, sum = 0;

    for (i = 0; i < M; i++)
        for (j = 0; j < N; j++)
            sum += a[i][j];
    return sum;
}
```

Resposta: Sim. Os compiladores de C guardam as matrizes, na memória, linha a linha. Neste caso o percurso na matriz é feito linha a linha. Portanto faz-se sempre acesso a posições de memória contíguas.

AC - 2016/2017

31

Localidade (2)

- Esta função tem boa localidade?

```
int sumarraycols(int a[M][N])
{
    int i, j, sum = 0;

    for (j = 0; j < N; j++)
        for (i = 0; i < M; i++)
            sum += a[i][j];
    return sum;
}
```

Resposta: Não. Neste caso o percurso da matriz é feito coluna a coluna; a sequência é $a[0][0]$, $a[1][0]$..., $a[M][0]$, $a[0][1]$... Portanto dois acessos consecutivos não fazem acesso a posições de memória contíguas.

AC - 2016/2017

32

Tecnologia de memória

- Todas as memórias rápidas são baseadas em semicondutores
- *Dynamic Random Access Memory (DRAM)*
 - DRAM oferece a melhor relação preço/desempenho, mas necessita de refrescamento periodicamente e após cada leitura
- *Static Random Access Memory (SRAM)*
 - SRAM é mais rápida mas muito mais cara, consome mais energia e ocupa mais espaço
- A memória RAM dos computadores é DRAM

AC - 2016/2017

33

Sumário SRAM vs DRAM

	Transistores por bit	Tempo acesso	Precisa refresh?	Custo	Utilização
SRAM	4 or 6	1X	Não	100x	Memória cache
DRAM	1	10X	Sim	1X	Memória central, frame buffers

- DRAM e SRAM são memória volátil
 - Perdem informação se a energia é desligada.

AC - 2016/2017

34

Hierarquia de memória

- Propriedades fundamentais:

- As tecnologias mais rápidas custam mais por byte e têm menor capacidade.
- O fosso entre o CPU e a memória central está a aumentar.
- Os bons programas têm boa localidade.

- Estas propriedades suportam:

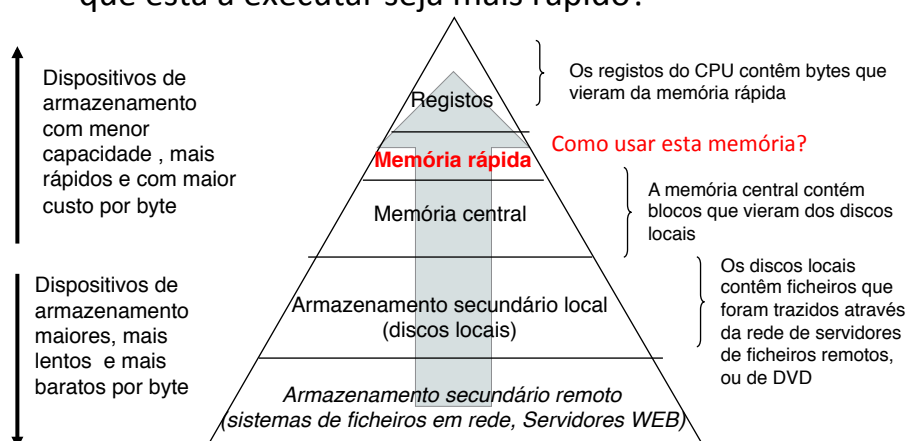
- abordagem do armazenamento como uma hierarquia de memória
- Introduzir mais um nível de memória para tirar partido da localidade: queremos memória central DRAM com o desempenho da SRAM

AC - 2016/2017

35

Hierarquia de níveis de memória

- Poderemos ter mais um nível para que o programa que está a executar seja mais rápido?



AC - 2016/2017

36

Ideia da *Cache*

- **Cache:** Um dispositivo “invisível” que atua como zona temporária para armazenamento de dados de outro dispositivo mais lento mas maior.
- Ideias fundamentais:
 - No nível L, o dispositivo menor e mais rápido, guarda parte do conteúdo do dispositivo maior e mais lento do nível L+1
 - Baseado na localidade, procura-se ter no nível L os dados de L+1 mais prováveis de ser brevemente usados
 - A atualização dos dados em L deve ser transparente para o utilizador desse nível

AC - 2016/2017

37

Utilização da ideia

- O mesmo princípio é explorado a outros níveis
- Exemplo dos *browsers Web*:
 - Procuram manter no disco local os dados remotos
 - Mantém em memória os dados que estão a usar ou que podem vir a usar brevemente

AC - 2016/2017

38

Ideia da Cache

● Vantagens esperadas:

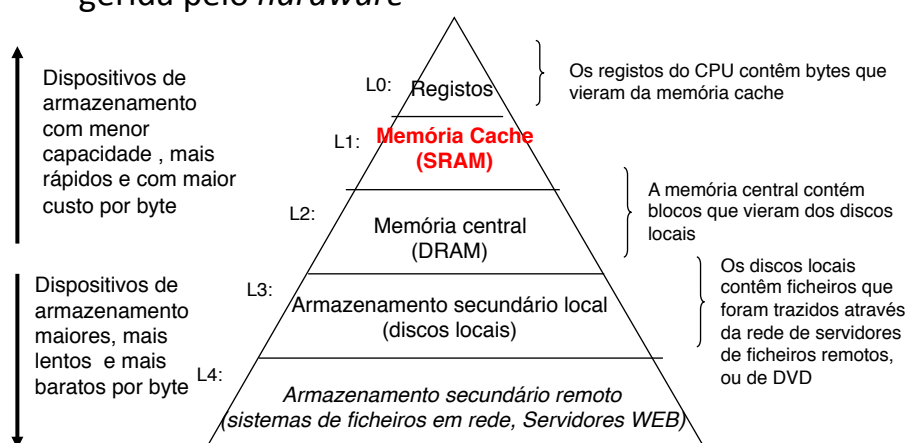
- A localidade leva a que a maior parte dos acessos (leitura/escrita) sejam aos dados já no nível L em vez de ter de ao nível L+1.
- O armazenamento no nível L+1 pode ser mais lento, e assim mais barato e muito maior.
- Efeito global: Um grande conjunto de memória que custa o preço da que se encontra no nível mais baixo, mas que permite o acesso aos dados a um ritmo semelhante à do nível mais elevado.

AC - 2016/2017

39

Hierarquia de níveis de memória

● Poderemos ter um, 2, ou 3 níveis de memória cache, gerida pelo *hardware*



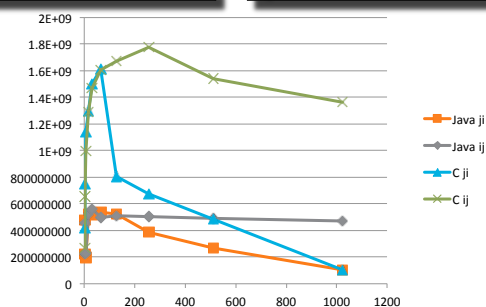
AC - 2016/2017

40

Desempenho (somas/s)

```
/* ij */
sum = 0.0;
for (i=0; i<n; i++) {
  for (j=0; j<n; j++) {
    sum += a[i][j];
  }
}
```

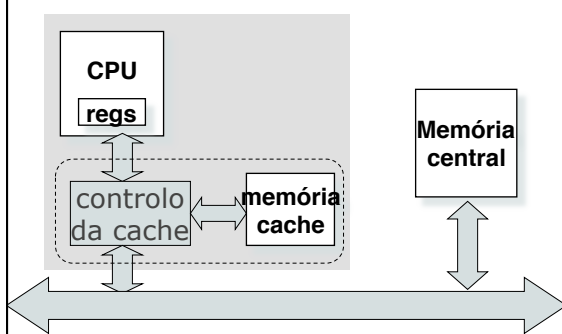
```
/* ji */
sum = 0.0;
for (j=0; j<n; j++) {
  for (i=0; i<n; i++) {
    sum += a[i][j];
  }
}
```



AC - 2016/2017

41

Organização da cache de memória



● Funcionamento:

CPU referencia memória:

se está em cache

accede à cache

senão

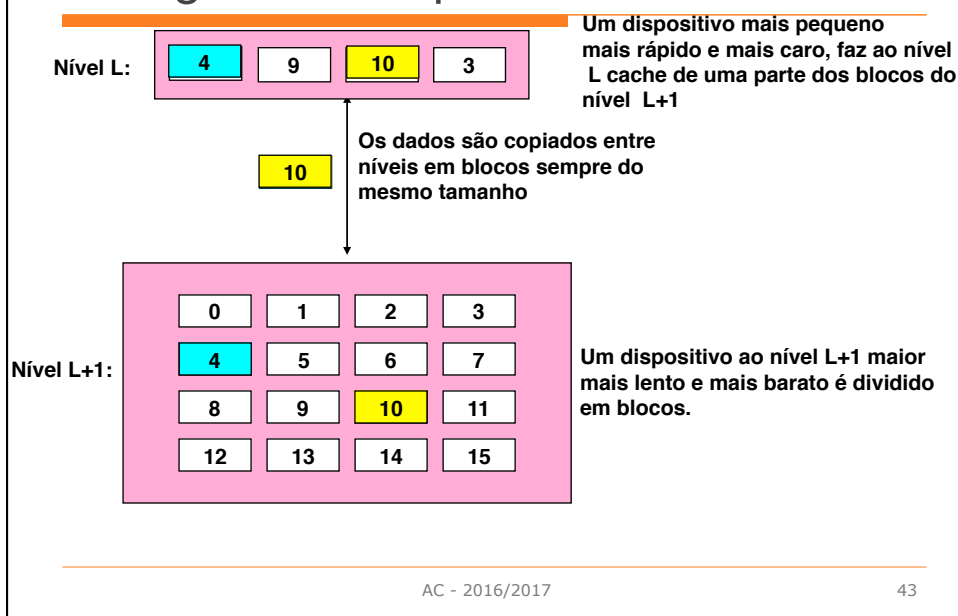
accede à memória

e actualiza cache

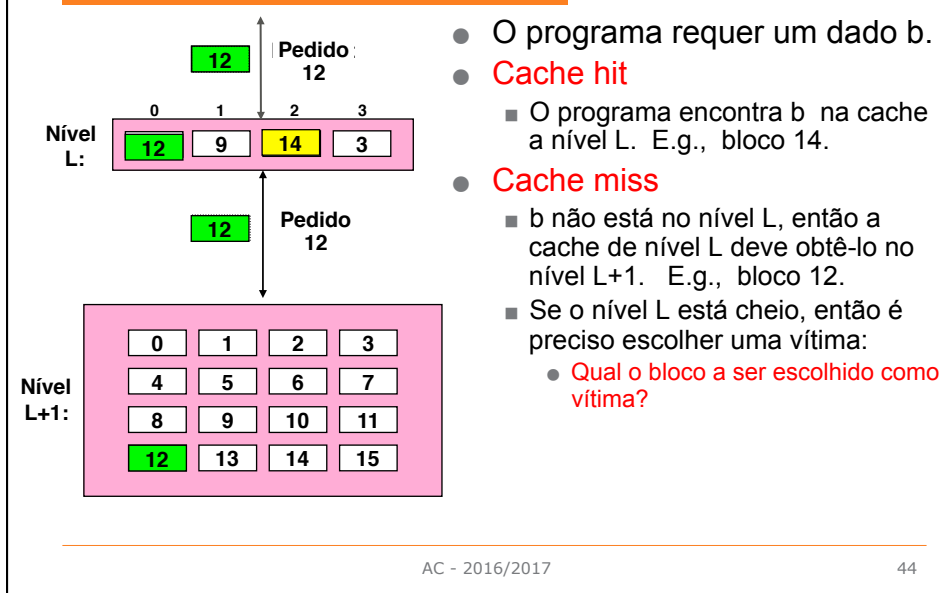
AC - 2016/2017

42

Caching na hierarquia de memória



Conceitos gerais da cache



Conceitos gerais da cache: hit/miss

- **Cache hit**
 - O programa encontra o dado pretendido na cache (nível L). Não é preciso fazer nada!
- **Cache miss**
 - Não encontra o que pretende no nível L, então a cache de nível L deve obtê-lo no nível L+1 (e assim sucessivamente).
 - Se o nível L está cheio, então é preciso arranjar espaço, escolhendo uma *vítima*:
 - Qual o bloco a ser escolhido como vítima (ao acaso? LRU - *least recently used* ?)
 - Se o bloco vítima está *limpo*, isto é, **não foi alterado** desde que veio para o nível L, posso carregar o novo bloco por cima
 - Se o bloco vítima está *sujo*, isto é, **foi alterado**, é preciso escrevê-lo no nível L+1 antes de o substituir

AC - 2016/2017

45

Conceitos gerais duma cache

- Taxa de sucesso (*hit ratio*)
 - h = percentagem de vezes que o dado pretendido está na cache:

$$h = \text{núm. cache hit} / \text{núm. total acessos}$$
- Tempo de acesso médio:
 - T_L é o tempo de acesso no nível L
 - T_{L+1} é o tempo de acesso no nível L+1
 - Tempo de acesso médio: T_a

$$T_a = h * T_L + (1 - h) * T_{L+1}$$

AC - 2016/2017

46

● Exemplo:

■ $T_{cache} = 2 \text{ ns}$

■ $T_{mem} = 8 \text{ ns}$

■ $Hit \text{ ratio} = 80\%$

● $T_a = h * T_L + (1 - h) * T_{L+1}$

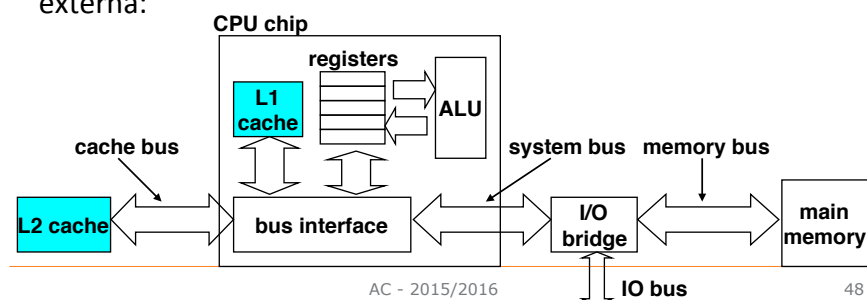
■ $T_a = 0,8 * 2 + 0,2 * 8 = 3,2 \text{ ns}$

AC - 2016/2017

47

Cache – Exemplo com dois níveis

- Cache são memórias rápidas (SRAM) geridas automaticamente pelo *hardware*.
- CPU procura em L1, depois em L2, finalmente na memória central
 - Podem introduzir algum *overhead*, mas desprezável
- Exemplo de arquitetura com L1 interna ao chip do CPU e L2 externa:



AC - 2015/2016

48

Hierarquia de níveis de memória

