

Arquitetura de Computadores

MiEI – 2016/17

DI-FCT/UNL

Aula 18

AC - 2016/2017

1

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.

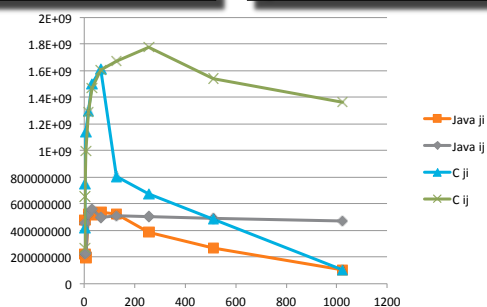
AC - 2016/2017

2

Desempenho (somas/s)

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

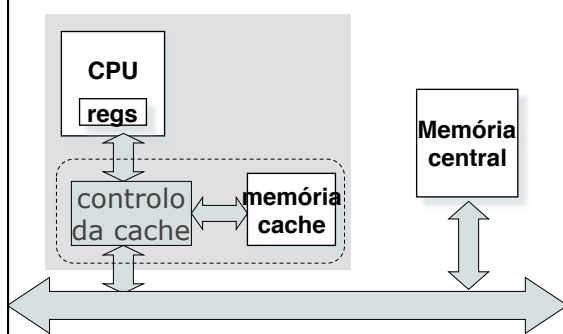
```
/* ji */
sum = 0.0;
for (j=0; j<n; j++) {
  for (i=0; i<n; i++) {
    sum += a[i][j];
  }
}
```



AC - 2016/2017

3

Organização da cache de memória



● Funcionamento:

CPU referencia memória:

se está em cache

accede à cache

senão

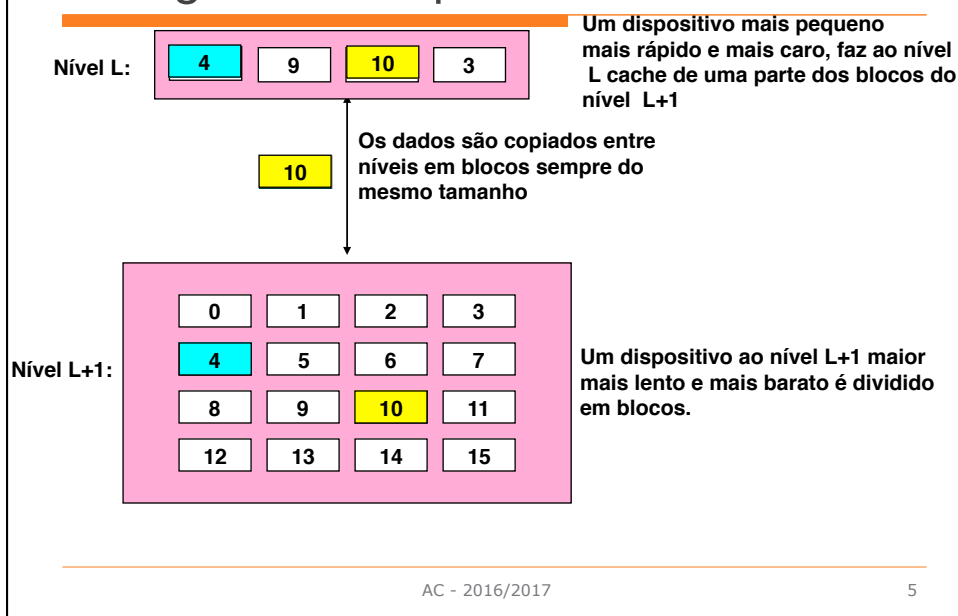
accede à memória

e actualiza cache

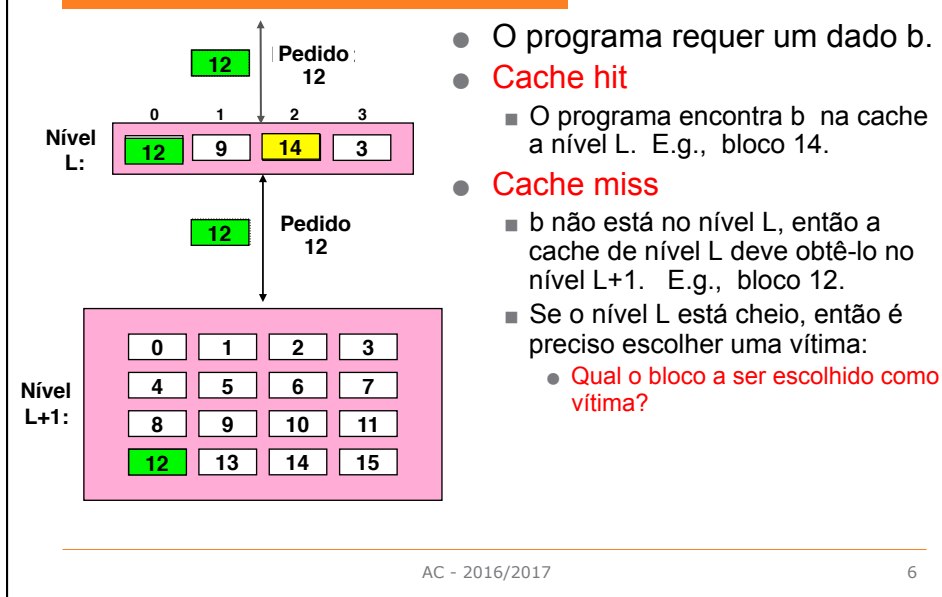
AC - 2016/2017

4

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 mais 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

7

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

8

● Exemplo:

■ $T_{cache} = 2 \text{ ns}$

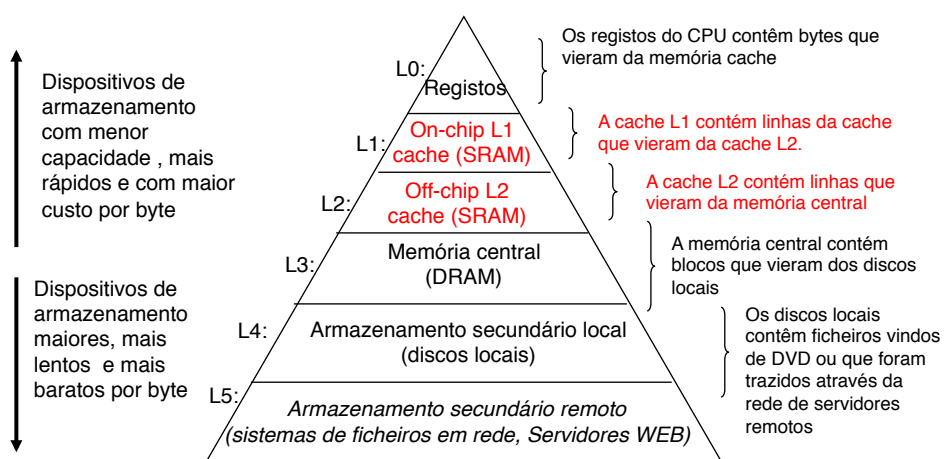
■ $T_{mem} = 8 \text{ ns}$

■ $Hit \text{ ratio} = 80\%$

● $T_a = h * T_L + (1 - h) * T_{L+1}$

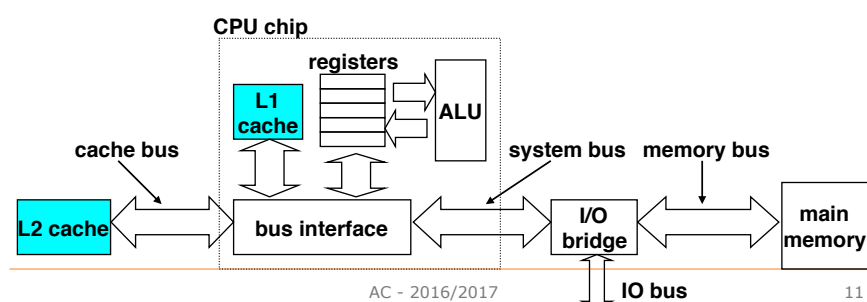
■ $T_a = 0,8 * 2 + 0,2 * 8 = \mathbf{3,2 \text{ ns}}$

Hierarquia de níveis de memória



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 L2 externa ao chip do CPU:



AC - 2016/2017

11

Política de escritas com cache

- Quando o CPU altera o conteúdo numa posição de memória, faz sentido que a alteração fique em cache
 - a ideia é que o CPU pode vir ler esse valor em breve (localidade temporal)
- O que acontece na memória de nível inferior, quando um bloco na cache é alterado, depende da **política de escrita**:
 - *Write-through*
 - *Write-back* (ou *delayed*)

AC - 2016/2017

12

Write-through

- A memória cache e a memória central são actualizadas em cada escrita
- Assegura que a cache e a memória central estão sempre **coerentes**
- Cada escrita demora o tempo de escrita na memória central
- Tolerável, porque normalmente há muito mais acessos em leitura do que em escrita

AC - 2016/2017

13

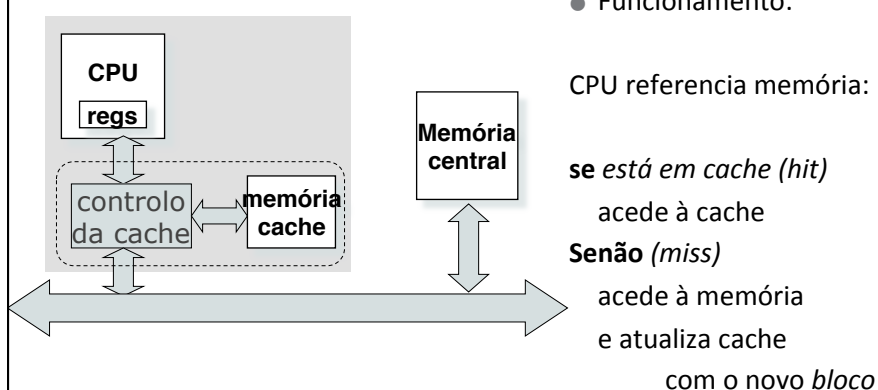
Write-back

- Quando há uma escrita só é actualizada a cache; mas marca-se esse bloco como alterado (*dirty*)
- Só quando um bloco é removido da cache é que se actualiza o bloco na memória de nível inferior
- Isto é mais rápido do que o *write-through*
- Diminui o tráfego para a memória de nível inferior
- O inconveniente é que a cache e a memória inferior nem sempre estão coerentes
 - isto pode provocar problemas ...

AC - 2016/2017

14

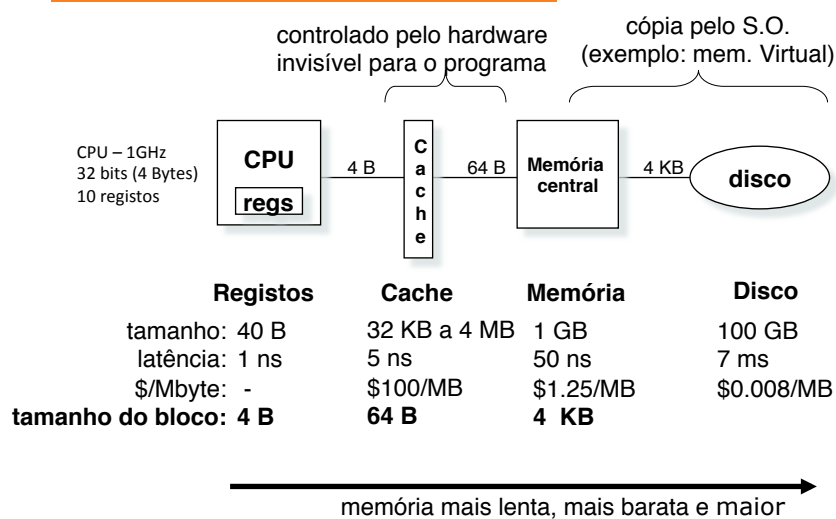
Organização da cache de memória



AC - 2016/2017

15

Exemplos de blocos usados

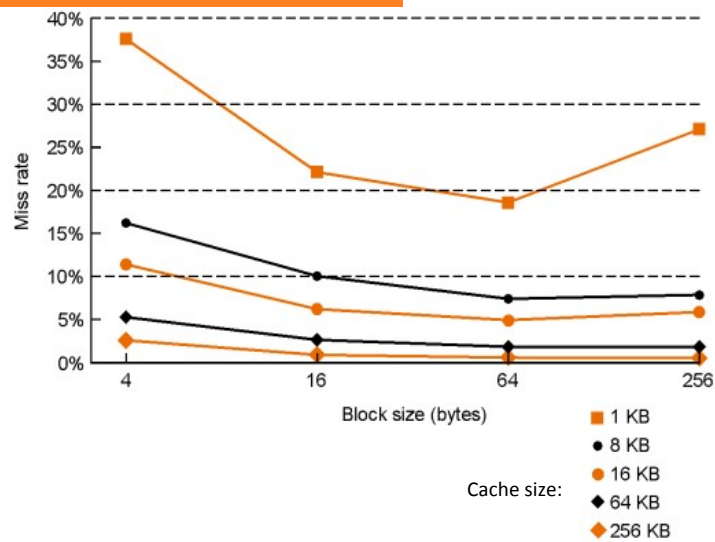


AC - 2016/2017

16

Impacto do tamanho do bloco no *hit-rate*

EXEMPLO
para um
programa:

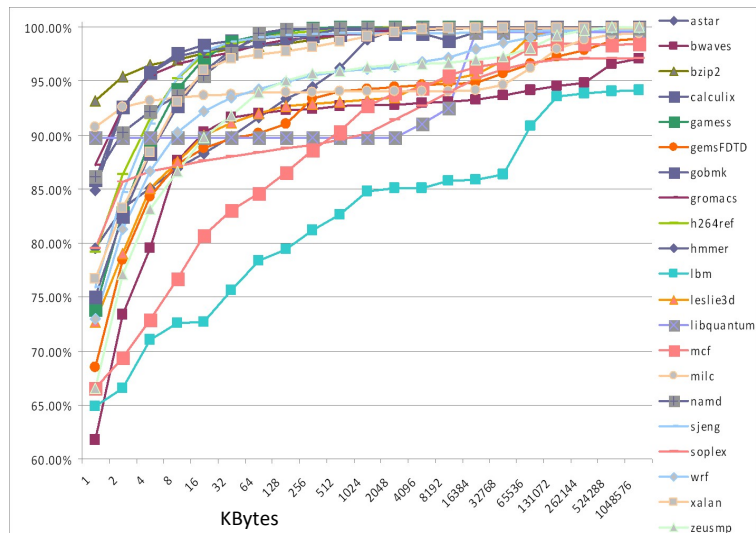


AC - 2016/2017

17

Impacto do tamanho da cache no *hit-rate*

EXEMPLOS
de vários
programas



AC - 2016/2017

18

Organização da cache

- Para além de guardar o conteúdo da memória de nível inferior, também necessita de guardar:
 - Que zonas da cache estão realmente em uso (de início a cache está vazia!)
 - Que endereços/blocos da memória de nível inferior estão na cache
 - Se usar *write-back*, que zonas foram escritas
 - Informação que permita escolher os blocos a substituir
- A consulta e atualização desta informação **tem de ser eficiente!!!**

AC - 2016/2017

19

Endereços e blocos de memória

Exemplo:

endereços de 8 bits

blocos de 4bytes

dado um endereço E :

$E/4$ = número do bloco

$E\%4$ = byte dentro do bloco

ou seja:

4 bytes por bloco = $2^2 \rightarrow 2$ bits

os 2 bits menos significativos indicam bytes dentro do mesmo bloco

os restantes correspondem ao número do bloco



núm. bloco deslocamento no bloco

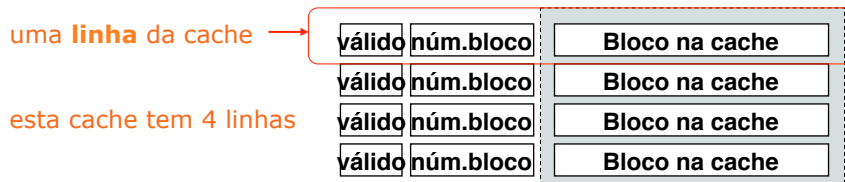
0=	00000000	
1=	00000001	bloco 0
2=	00000010	
3=	00000011	
4=	00000100	bloco 1
5=	00000101	
6=	00000110	
7=	00000111	
8=	00001000	bloco 2
9=	00001001	
10=	00001010	
11=	00001011	
12=	00001100	bloco 3
13=	00001101	
14=	00001110	
15=	00001111	
16=	00010000	bloco 4
17=	00010001	
18=	00010010	
19=	00010011	
20=	00010100	

AC - 2016/2017

20

Exemplo (*write-through*)

- A cache tem para cada bloco, se está ocupado (válido) ou não e o respectivo número de bloco, quando ocupado



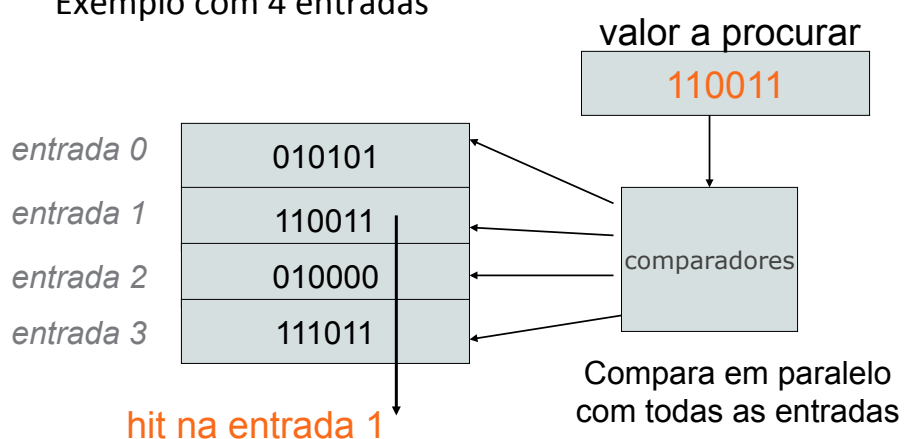
- Para cada acesso do CPU à memória, parte-se o endereço de acordo com o tamanho do bloco e procura-se o seu número na cache
 - só faz sentido se esta procura for rápida!
- Usa-se os bits de deslocamento no bloco para aceder ao byte pretendido

AC - 2016/2017

21

Memória associativa (interrogável pelo conteúdo)

Exemplo com 4 entradas

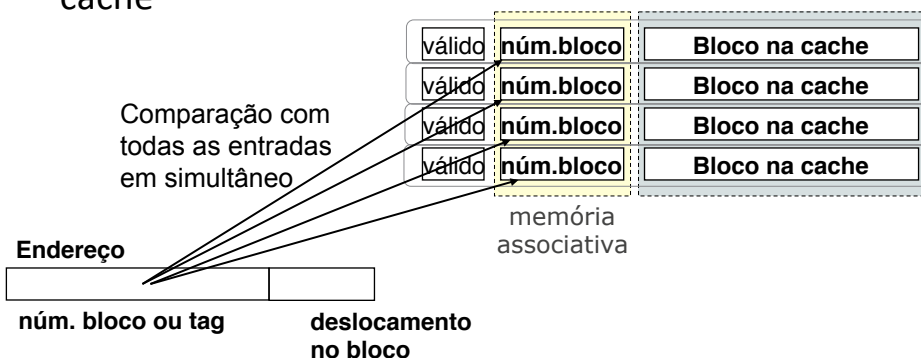


AC - 2016/2017

22

Cache associativa pura

- Usa o número do bloco como **chave** (ou *tag*) na memória associativa para procurar o bloco na cache

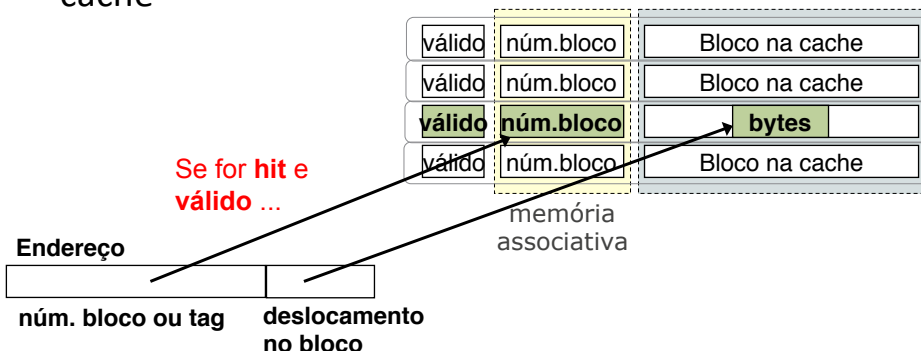


AC - 2016/2017

23

Cache associativa pura

- Usa o número do bloco como **chave** (ou *tag*) na memória associativa para procurar o bloco na cache



AC - 2016/2017

24

Caches associativas puras

- As Caches Associativas são muito eficientes:
 - qualquer bloco pode estar em qualquer posição
 - a pesquisa do bloco faz-se em paralelo, com tantos comparadores quantos os blocos que cabem na cache
- Mas são complexas e muito caras:
 - cada linha da cache deve guardar todos os bits do número do bloco nessa posição
 - cada comparador é de tantos bits quantos os bits no número do bloco
 - muitos comparadores, um para cada linha da cache

AC - 2016/2017

25

Tentando não usar memória associativa

- Fixamos a posição na cache para cada bloco
- Exemplo:
 - cache com 4 linhas de blocos de 4 bytes
 - o núm. da linha é dado pelo $\text{núm. bloco} \% \text{núm. linhas}$
 - ou seja, por 2 bits do endereço:



linha 0	Bloco 0, 4, 8, ...
linha 1	Bloco 1, 5, 9, ...
linha 2	Bloco 2, 6, 10, ...
linha 3	Bloco 3, 7, 11, ...

```

0= 00000000
1= 00000001
2= 00000010 bloco 0
3= 00000011
4= 00000100
5= 00000101
6= 00000110 bloco 1
7= 00000111
8= 00010000
9= 00010001
10= 00010010 bloco 2
11= 00010011
12= 00010100
13= 00010101 bloco 3
14= 00010110
15= 00010111
16= 00011000
17= 00011001 bloco 4
18= 00011010
19= 00011011
20= 00011100 bloco 5

```

AC - 2016/2017

26

Porque se usam os bits do meio para escolher a linha?

Cache com 4 linhas

00	
01	
10	
11	

- Se usasse os bits mais signif.:
 - Blocos de memória adjacentes corresponderiam à mesma linha na cache
 - Mau para a localidade espacial !
- Usando os bits do meio:
 - Blocos consecutivos de memória correspondem a linhas da cache distintas
 - Melhor para a localidade espacial !

Utilizar os bits Mais significativos

0000	
0001	
0010	
0011	
0100	
0101	
0110	
0111	
1000	
1001	
1010	
1011	
1100	
1101	
1110	
1111	

Utilizar os bits do meio

0000	
0001	
0010	
0011	
0100	
0101	
0110	
0111	
1000	
1001	
1010	
1011	
1100	
1101	
1110	
1111	

AC - 2016/2017 27

Cache de mapa directo

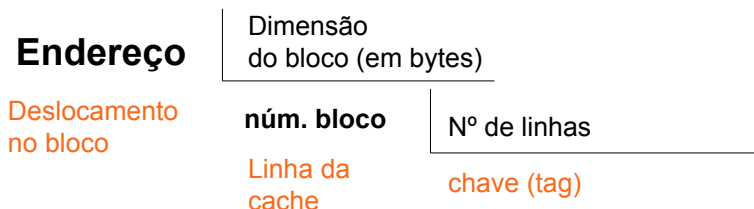
- É necessário manter para cada linha uma chave (ou *tag*) que identifique qual o bloco realmente na cache de entre os vários possíveis:
 - o que distingue os vários blocos são os bits mais significativos restantes do endereço
- Exemplo: end. de 8bits, cache de 4 linhas e 4 bytes p/bloco:

chave núm. linha desl. no bloco
 Endereço:

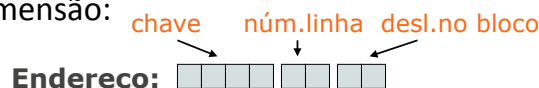
linha 0	válido	chave	Bloco na cache
linha 1	válido	chave	Bloco na cache
linha 2	válido	chave	Bloco na cache
linha 3	válido	chave	Bloco na cache

Tratamento do endereço

● Resumo (divisões inteiras):



- em binário, estas operações correspondem a separar os bits do endereço em 3 partes, de acordo com cada dimensão:



AC - 2016/2017

29

Vantagens da Cache de mapa directo

- As Caches de mapa direto são mais baratas do que as Associativas:
 - O mapa da Cache tem menos bits;
 - Não necessita de memória associativa e só precisa de 1 comparador (de tantos bits quantos os na chave dos blocos)
- São eficientes na pesquisa em Cache:
 - porque usam os bits do endereço diretamente como índice no mapa da Cache e como chave, só precisando de comparar essa chave.

AC - 2016/2017

30

Desvantagens da Cache de mapa directo

- Levam a colisões: todos os blocos com a mesma linha atribuída colidem nessa posição
 - mesmo que existam outras linhas na cache livres!
- A probabilidade de colisões diminui à medida que se aumenta a Capacidade da Cache
 - mas o problema mantém-se ...