

Lezione 18

Introduzione alla memoria cache

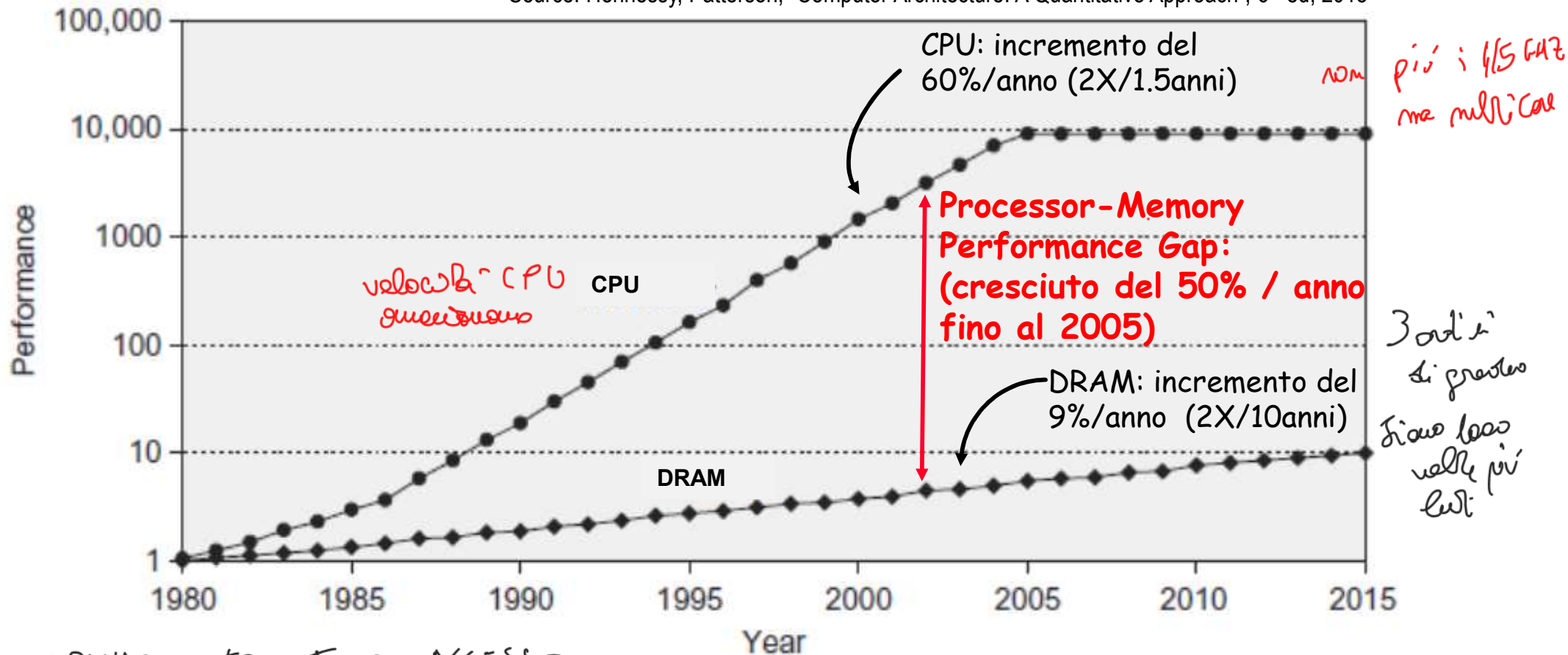
cache, funzione extra per

Patterson: 5.3

migliorare prestaz. macchina

• Gap Prestazioni Processore-Memoria

Source: Hennessy, Patterson, "Computer Architecture: A Quantitative Approach", 6th ed, 2018



GRANDE DIVARIO tra Tempo ACCESSO

alla MEMORIA e Tempo con cui

→ Necessità di nuove architetture per le memorie

Cui CPU vorrebbe SPERARE... sempre maggiore nel tempo

*velocità memoria non è aumentata molto,
siamo sulle stesse percentuali*

Gap Processore-Memoria: un esempio

Valutiamo il gap Processore-Memoria in termini di istruzioni per eseguite dal processore durante 1 accesso in memoria

Intel Core i7-9900k: $26 \text{ ns} / 0.2 \text{ ns} = 130 \text{ clk}$ $\times 4 \text{ istr/clk} \rightarrow 520 \text{ istr.}$

AMD Ryzen 2700X: $26 \text{ ns} / 0.23 \text{ ns} = 113 \text{ clk}$ $\times 4 \text{ istr/clk} \rightarrow 452 \text{ istr.}$

IBM Power8: $26 \text{ ns} / 0.3 \text{ ns} = 86.7 \text{ clk}$ $\times 8 \text{ istr/clk} \rightarrow 693 \text{ istr.}$

Tempo di accesso
alla memoria
(best-case DDR4)

Periodo di clock
del processore
(5GHz, 4.3GHz, 3.3GHz risp.)

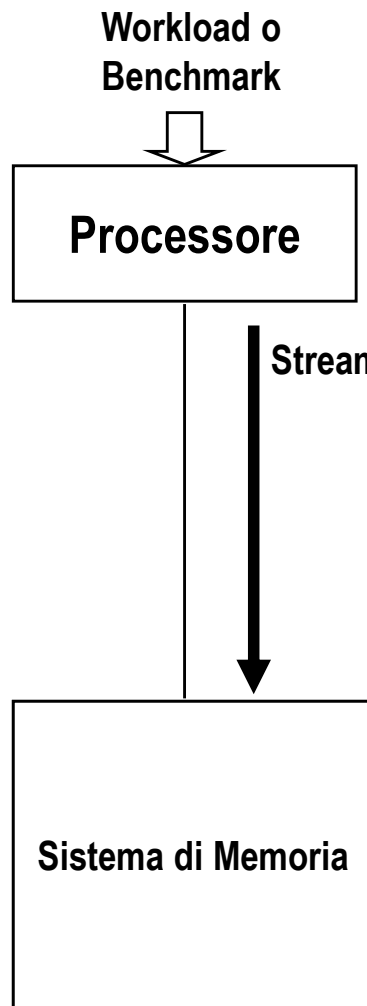
Istruzioni eseguite
in un ciclo di clock

**Istruzioni eseguite
dal processore
nel tempo in cui avviene
1 accesso in memoria**

*nel tempo in cui
facciamo un accesso alla
memoria, potrei
farne 500 istr.*

*⇒ Voglio ridurre gli accessi alla memoria
se non posso?*

"L'Arte della Progettazione del Sistema di Memoria"



osservo

Cosa accade sul BUS che collega Processore e Memoria

↳ Vedremo passare una serie di indirizzi
(dove fare fetch delle istruzioni) (oppure dove leggere e scrivere dalla Memoria Dati)

<op,addr>, <op,addr>, <op,addr>, <op,addr>, ...

Op = i-fetch, read, write
Addr = 0x00F0....

Op = i-fetch, read, write

Addr = 0x00F0....

qual'è cosa accade

Sul bus che collega CPU

Memoria... vedrai una serie di
indirizzi su cui fare Fetch

Possiamo ottimizzare il Sistema di Memoria in modo tale da minimizzare il tempo medio di accesso per i tipici programmi che utilizziamo (workload)?

Voglio un'architettura che faciliti gli accessi
al SISTEMA di MEMORIA

Principio di Località dei Programmi Proprietà fondamentale

Se osservo gli indirizzi accesi

- Un riferimento in memoria tende ad essere ripetuto dopo poco tempo

→ LOCALITA' TEMPORALE

Se osservo gli indirizzi che la CPU genera....

- Esempio:

- 0008
- 1000
- 0008
- 2000
- 0008

l'indirizzo 0 contiene una variabile del programma,

IN UN BREVE PERIODO HO ACCESSI CONTINUI ALLA STESSA LOCALITÀ DI MEMORIA

- Un riferimento in memoria tende ad essere a locazioni vicine a quelle usate recentemente

→ LOCALITA' SPAZIALE

es. Fetch del codice, le istruz. successive

- Esempio: = accesso ad indirizzi tra loro vicini

accesso ad instr. 4, 8, ...

- 0004
- 0008
- 000C
- 0010
- 0014

Se ho queste località,

posso allora derivare una architettura diversa del sist. di MEMORIA

⇒ Se ho queste località, la ripetizione di indirizzi, penso di mettere questi indirizzi e contenuto in una memoria più vicina al processore, CACHE

Implicazioni del Principio di Località

- 1) Posso utilizzare piccole memorie veloci per mantenere i riferimenti in memoria più utilizzati dal processore

(indivisi e contenuti)

→ **MEMORIA CACHE**

piccola memoria che

tiene i dati più recenti che sto utilizzando in memoria

↳ così posso fare accesso in maniera più rapida / uso SRAM

- 2) Se i riferimenti utilizzati non possono essere soddisfatti nella memoria veloce

Se non trovo dati nella cache, faccio accesso alla memoria principale

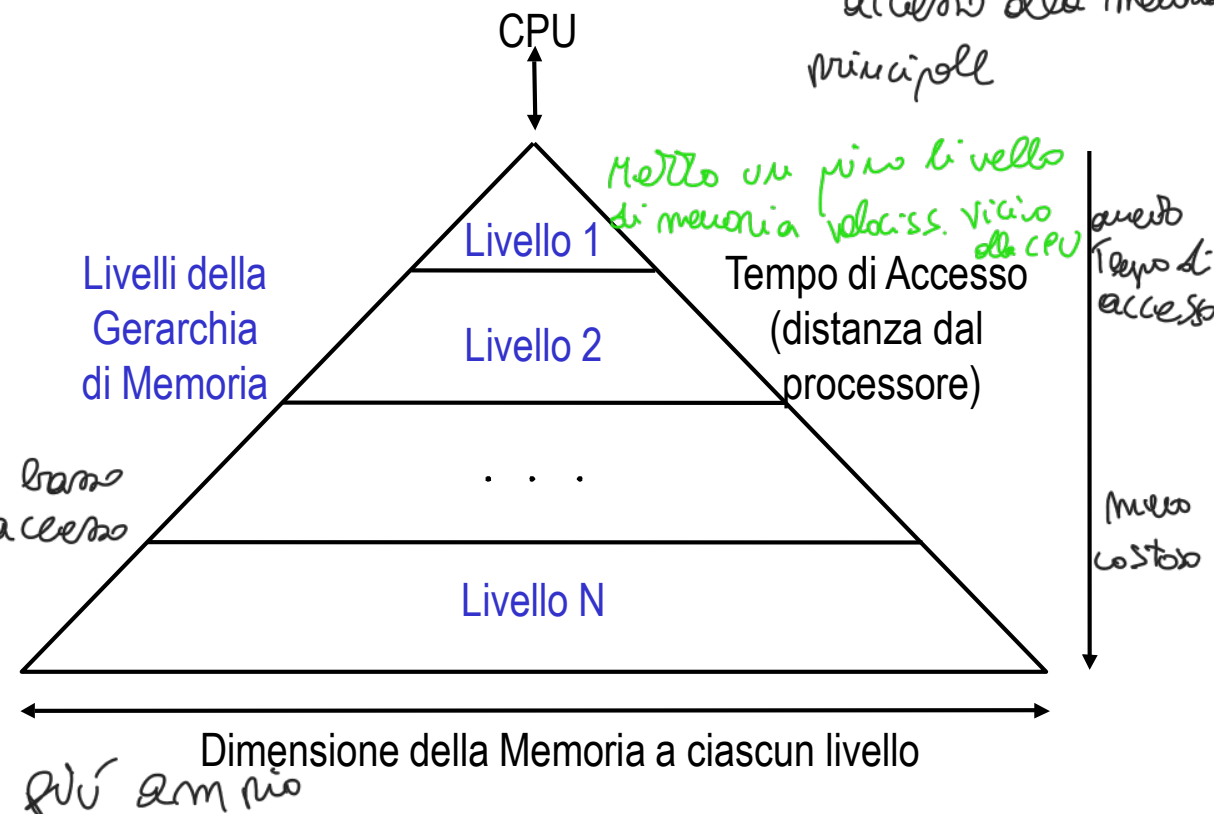
→ creo un secondo livello di memoria un po' più grande *più ampio* e un po' meno veloce

- 3) Posso reiterare questo ragionamento per N livelli di memoria

avuto una

→ **GERARCHIA DI MEMORIA**

⇒ andando verso il basso aumento i tempi di accesso e le dimensioni



con gerarchia, altro possibilità di trovare dato in un certo livello, hit o miss se dato non disponibile

Gerarchia di Memoria

- Presenta al processore una quantità di memoria pari a tutta quella disponibile all'ultimo livello (realizzata con la tecnologia più economica)
 lo cerco nel liv. succ. e' conveniente lavorare trasferendo 32-64 BYTE es. blocchi
- Fornisce un tempo di accesso che si avvicina a quello del primo livello (in cui posso usare la tecnologia più veloce possibile). Es. con 2 soli livelli:
 - Primo Livello: Cache, tecnologia SRAM
 - Secondo Livello: Memoria Principale, tecnologia DRAM
 voglio dare l'illusione di fornire processore la quantità di memoria dell'ultimo livello con tempi di accesso del primo

• Matematicamente:

- h = probabilità di trovare un dato nel primo livello - *Maggiore LOCALITA' comporta $h \rightarrow 1$ HIT*
- m = probabilità di trovare un dato nel livello successivo (ovviamente $h+m=1$) *MISS*
- t_{hit} = tempo di accesso al dato dal primo livello *es. 2ns*
- t_{miss} = tempo di accesso al dato nel livello successivo *200ns*

→ Average Memory Access Time (AMAT):

$$AMAT = h \times t_{hit} + m \times t_{miss}$$

se $h \rightarrow 1$ allora $AMAT \rightarrow t_{hit}$ massima velocità

Se scomponiamo MissTime = HitTime + PenaltyTime *Tempo aggiuntivo per andare al 2°*

$$= h \times t_{hit} + m \times (t_{hit} + t_{penalty}) = t_{hit} + m \times t_{penalty}$$

penalty, andare al 2° lev.

ho sicurezza in tempo di accesso che è minimo il tempo di hit, t, eventuale perdita

Terminologia

blocco: insieme di byte a locazioni contigue ed indirizzi multipli della dimensione del blocco

hit: il riferimento può essere soddisfatto nel livello di memoria immediatamente successivo

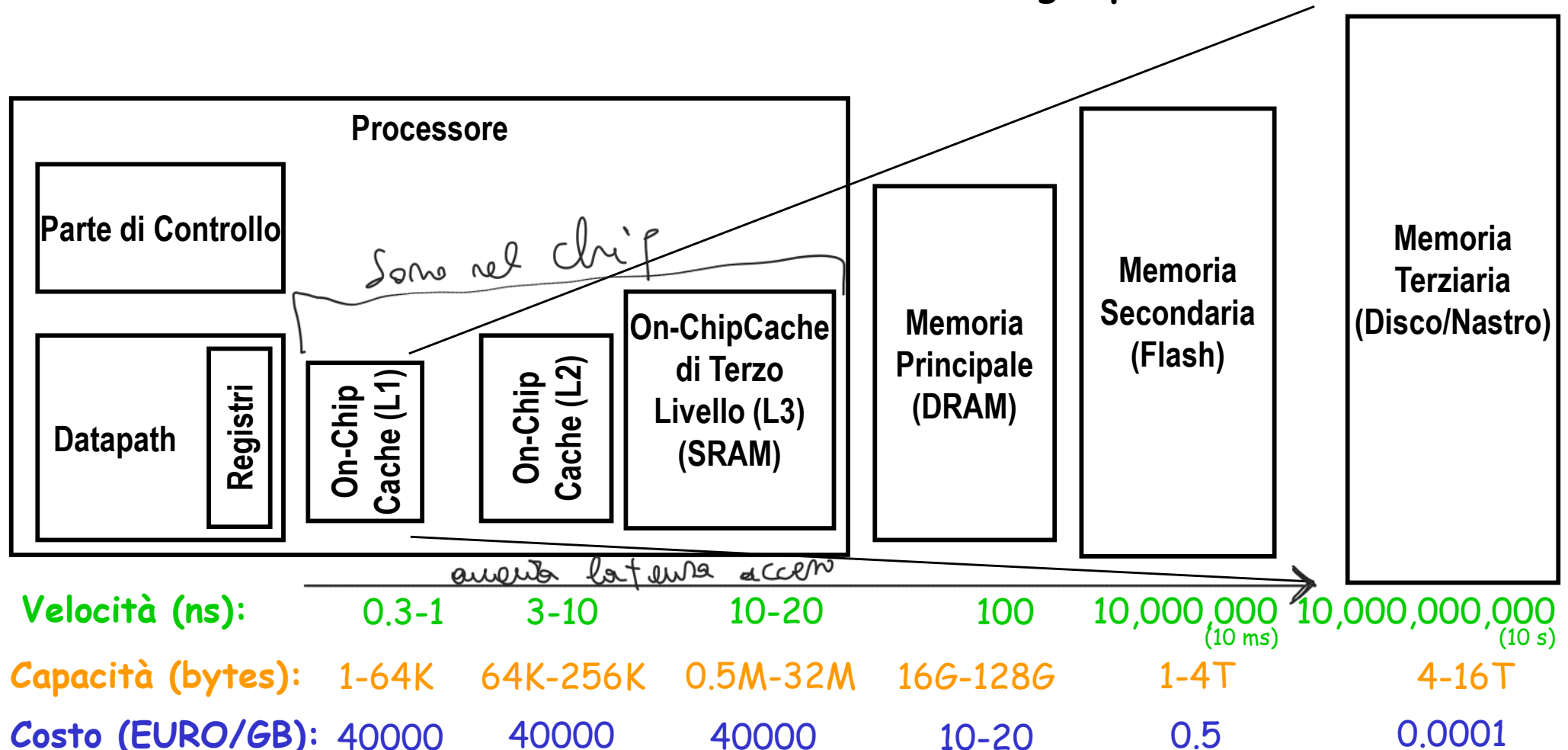
miss: il riferimento richiede un accesso a livelli di memoria oltre il successivo (più lenti) *dato non disp. in quel livello*

$$AMAT = T_{hit} + t_{miss} \cdot T_{penalty}$$

Transfer in un blocco di 64 BYTE non 64

Organizzazione gerarchica della memoria

- Sfruttando il principio di località è possibile:
 - Presentare all'utente tutta la memoria presente al livello della tecnologia meno costosa (es. disco o nastro)
 - Fornire la velocità di accesso offerta dalla tecnologia più veloce



Memoria Cache

La cache ragiona a blocchi, cerca di catturare il WORKING SET, i blocchi di Memoria più usati dal processore

- Scopo della **cache** è di memorizzare (per un rapido accesso) **copie di blocchi** di memoria principale *voglio catturare i dati che uso più freq.*
 - Funzionamento: ogni volta che faccio accesso ad una locazione di memoria, lascio in cache una copia del corrispondente blocco
 - Conseguenza: per il principio di località la cache tende a memorizzare i blocchi più usati dal processore (working set)

• Normalmente si ha:

- **B** = dimensione del Blocco di cache in byte = 2^b *saranno eguali in potenza del due*
- **C** = dimensione della Cache in byte = 2^c
- **M** = dimensione della Memoria in byte = 2^m *memoria principale*

$M \gg C$ per ragioni di costo

ovvero

$$N_M \gg N_C$$

N_M

= Numero di blocchi della Memoria = M/B

N_C

= Numero di blocchi della Cache = C/B

Inoltre, ad un certo istante per lo working set: *usato dal processore*

$W_C \subset W_M$ *Working Set (blocchi più usati del prog.) in Cache*

W_C = insieme dei blocchi del mio programma presenti in cache *saranno sempre un sottoinsieme di quelli nella RAM*

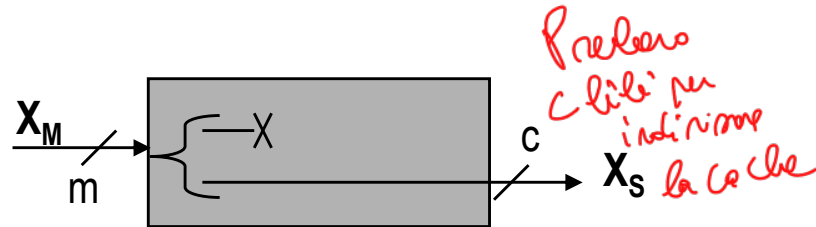
W_M = insieme dei blocchi (totali) del mio programma in memoria

ovvero tutto il working set in Cache, sempre HIT...

Cache ad ACCESSO DIRETTO

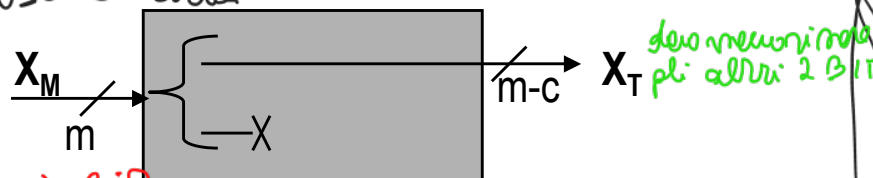
1) Parte dall'indirizzo specifico del processore al blocco X_M , prendo 3 bit meno significativi X_S per indirizzare la cache

- X_M = INDIRIZZO al BLOCCO
- X_S = BIT con cui INDIRIZZO la CACHE, *subindice*



Equivale a fare $X_S = X_M \bmod N_C$

Per capire quale dei 4 blocchi e' stato memorizzato in cache, uso etichetta



Equivale a fare $X_T = X_M \div N_C$

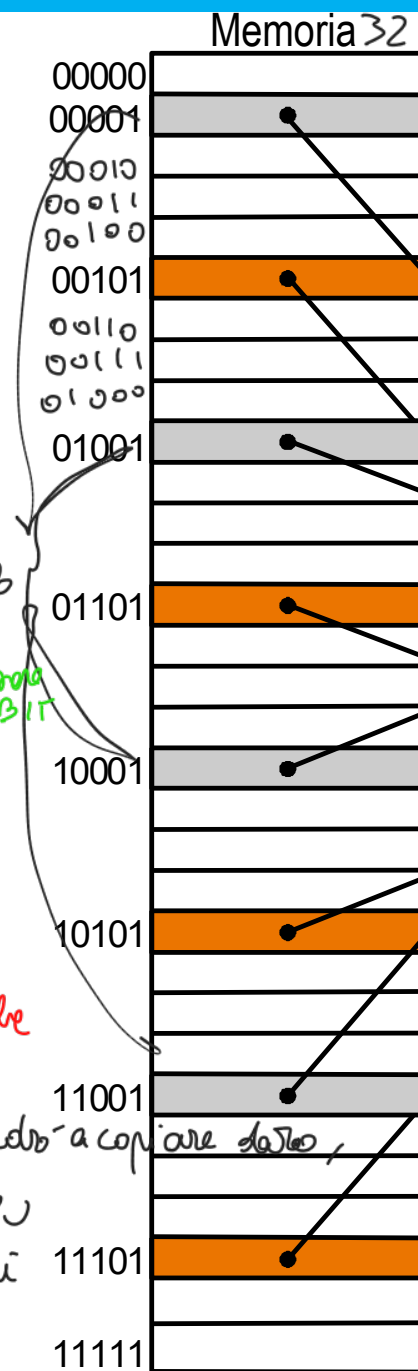
Prendo i bit
divisione rispetto al # blocchi di cache

X_S andrà a specificare lo slot di cache in cui andrò a copiare dato,

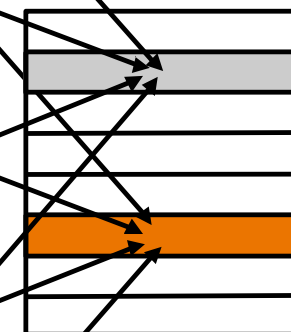
X_T servirà per confrontare se la CPU ha specificato uno di quelli 4 termini possibili

Esempio:
 $N_C = 8$, $N_M = 32$
($c = 3$, $m = 5$)

$\rightarrow X_S = F_S(X_M)$
 $= X_M \bmod N_C$
 $= X_M \bmod 8$



Cache 8 Blocchi



in quella slot di cache finiscono tutti i blocchi che finiscono con 001

in questa locazione di cache finiscono uno

$X_T = 00, 01, 10, 11$ finire uno

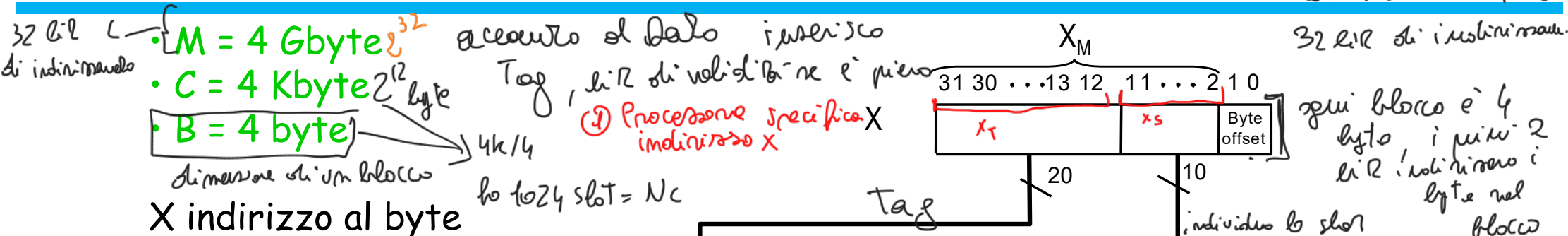
X_{Tag} bit di questi blocchi

residui dell'indirizzo
 X_T mi servono per ritrovare il blocco

Schema logico di una cache ad accesso diretto

es. Memoria a 4 Gb

32 bit di indirizzamento



indirizzo del blocco: $X_M = X / B = \text{indirizzo} / \text{dim. blocco}$

$= X / 2^2$

$N_c = \# \text{ slot} = C/B = 2^{10}$

to bit mi individuano lo slot

$X_S = X_M \bmod N_c \Rightarrow X_S = X_M \% \# \text{ slot}$

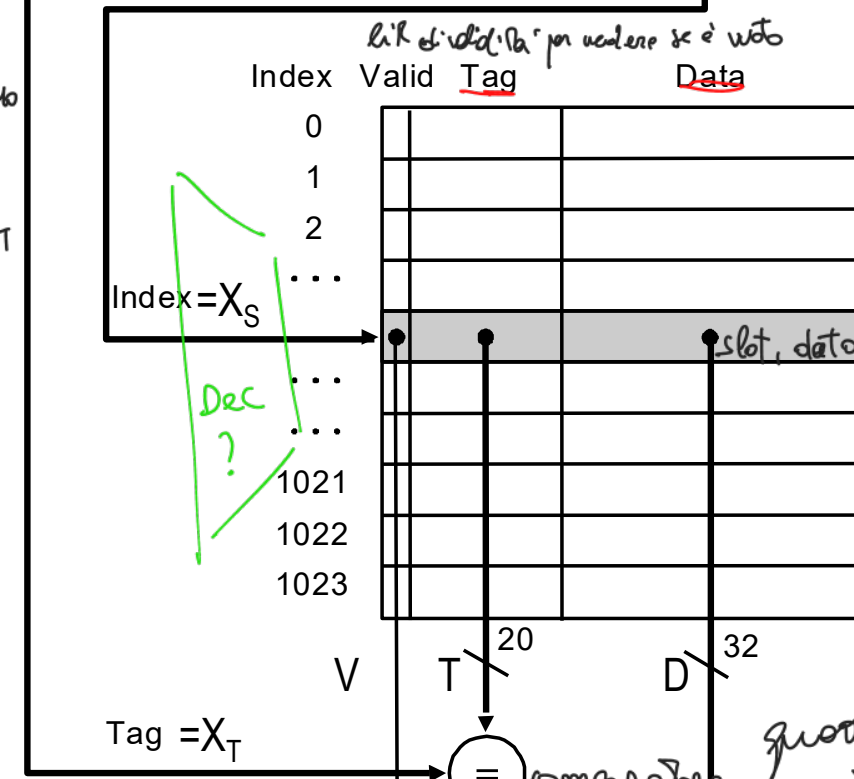
$= X_M \bmod C/B$

$= X_M \bmod 2^{10}$ prende l'indirizzo di

$X_T = X_M / N_c$

$= X_M / (C/B)$ 20 bit per il Tag

$= X_M / 2^{10}$



vero e proprio
dobbiamo essere sicuri
che il nostro dato
sia quello che lo
fatto qui

ho 2^{10} slot.

• NON APPENA LA CPU SPECIFICA UN INDIRIZZO MI VIENE IMMEDIATAMENTE DISPONIBILE QUESTA RIGA DI CACHE, con 3 segnali Valid, Tag, Data.

Se bit Valid = 1 e Tag corrisponde ($X_T = X_M / N_c$) allora lo bit, il dato lo prendo dalla cache e non dalla memoria principale

comparatore. Se bit valido, il dato è presente e non devo cercarlo in memoria

comparatore confronta bit tag con quello indirizzo

H $\rightarrow$ Hit

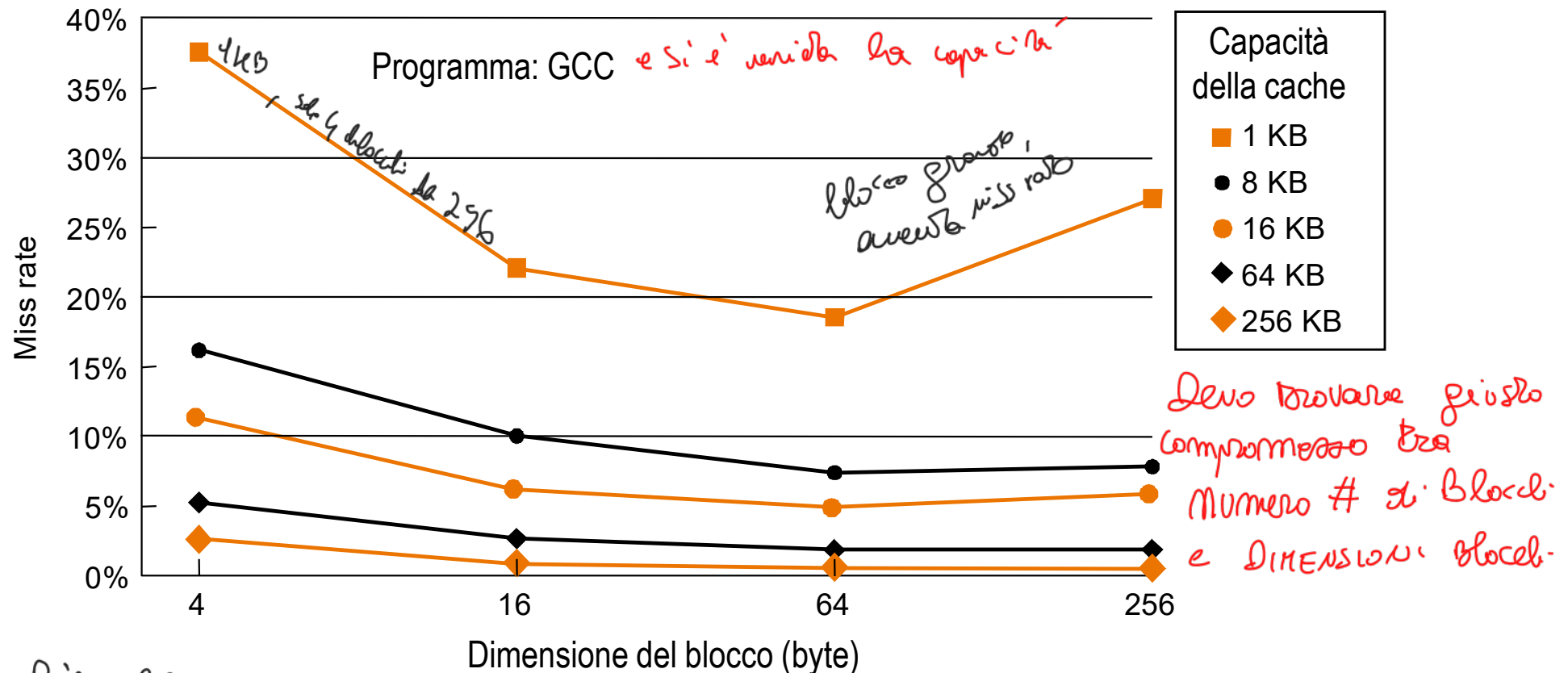
CPU prende il dato dalla cache se hit

quindi CPU specifica indirizzo, viene reso disponibile i 4 byte della cache a quell'indirizzo

Dipendenza del Miss Rate dalla dimensione del blocco

- Aumentando la dimensione del blocco il miss rate tende a diminuire

Meglio avere blocchi grandi o blocchi piccoli per migliorare # hit?



Program	Block size	Miss rate
gcc	4	5.4%
	16	1.9%
spice	4	1.2%
	16	0.4%

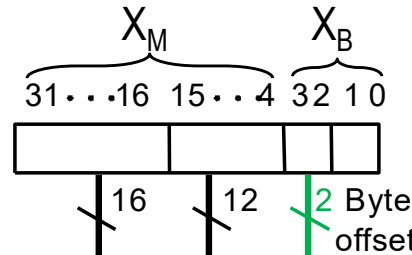
conveniente avere blocchi più grandi

Cache a blocchi più grandi: sfrutto la località spaziale

- $M = 4 \text{ Gbyte}$
- $C = 64 \text{ Kbyte}$
- $B = 16 \text{ byte}$

ogni riga della cache contiene 4
word da 4 byte, 128 bit

X indirizzo
al byte



$X_B \equiv X \bmod B$
 $X_B \equiv \text{byte Block-offset}$
 $X_O \equiv X_B / W$, e preferire la
 $X_O \equiv \text{word Block-offset}$
 $W \equiv \# \text{ di byte nella word}$

per selezionare una word
nel blocco per
il blocco per
e preferire la
word e i 32
bit
 corrisp.
sulla word cercata

Block offset = X_O

$$X_M = X / B$$

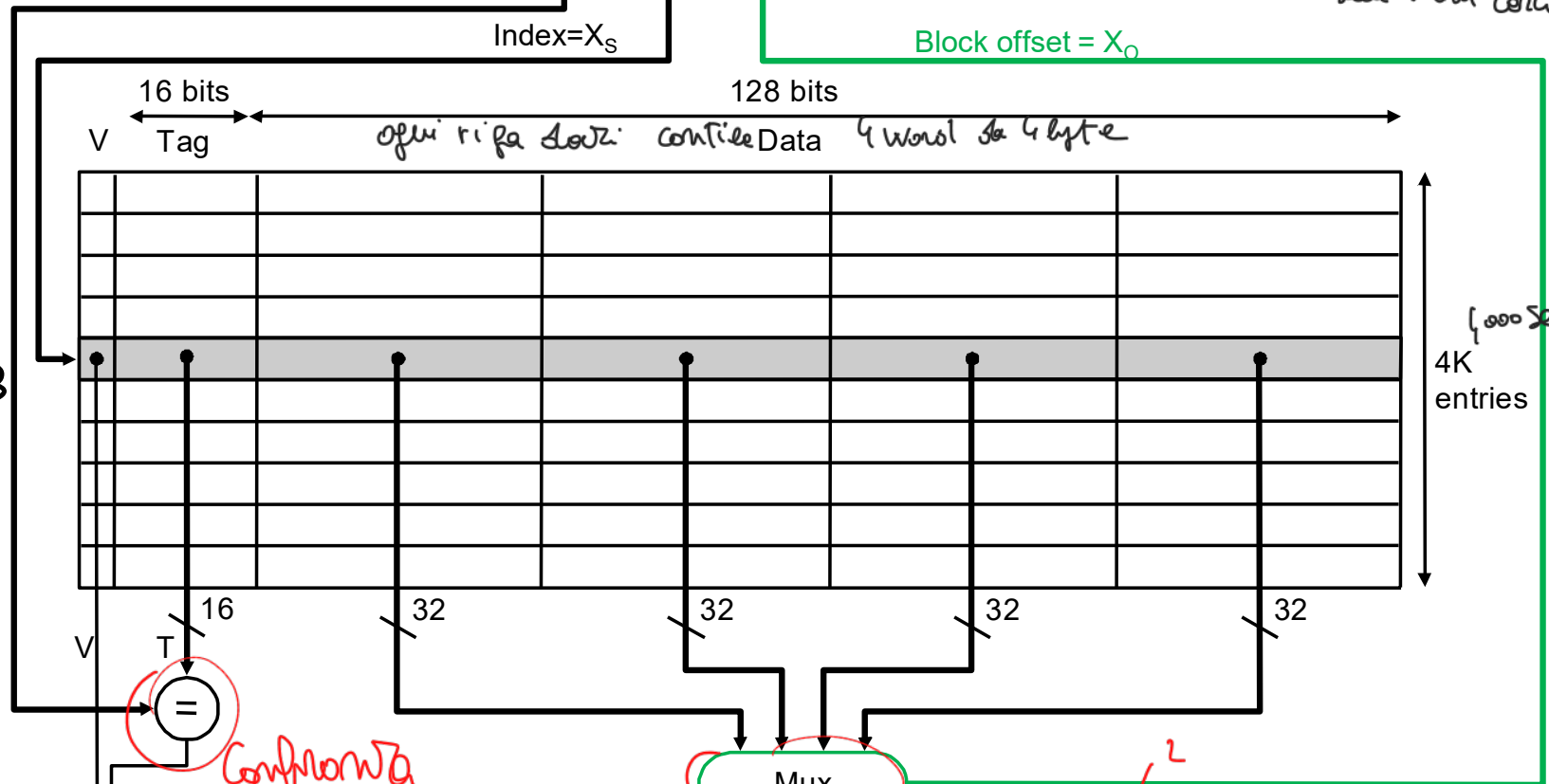
$$= X / 2^4$$

$X_S = 12$, bit per individuare
il set

$$X_S = X_M \bmod N_C$$

$$= X_M \bmod C/B$$

$$= X_M \bmod 2^{12}$$



4K
entries

16 bit per i tag

$$X_T = X_M / N_C$$

$$= X_M / (C/B)$$

$$= X_M / 2^{12}$$

Confronto
il tag con
quello dell'indirizzo

MUX

Mux

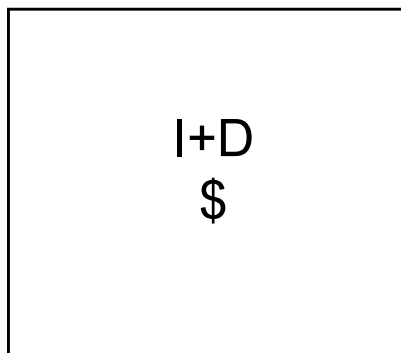
Per poi selezionare una word nel blocco
serve un multiplexer pilotato dai 2
bit dell'offset corrispondenti alla word
cercata

Split Cache

• Voglio migliorare ancora la cache

in realtà Memoria istruzione e memoria dati
non saranno altro che il primo livello di cache

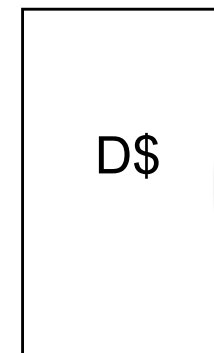
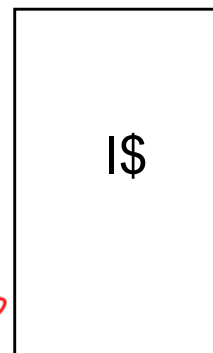
CACHE UNIFICATA L1 = Mem Istruz. e Memoria Dati



istruzioni
e dati

meglio split cache
poiché posso fare accessi
ad entrambe in parallelo
e miss rate migliore

SPLIT CACHE conviene:



posso fare accessi in parallelo

due versioni
quanto faccio
un Fetch nell'altro
• conserva meglio
la località

Accesso in
parallelo

le divide

- La split cache consente di sfruttare maggiormente la località spaziale soprattutto nell'accesso ai dati

• miss rate
migliore

Program	Block size in words	Instruction miss rate	Data miss rate	Effective combined miss rate
gcc	1	6.1%	2.1%	5.4%
	4	2.0%	1.7%	1.9%
spice	1	1.2%	1.3%	1.2%
	4	0.3%	0.6%	0.4%

Numero medio di riferimenti per istruzione

col ogni istruzione, la cache viene interrogata

facciamo FETCH dell'istruzione ADD, genera un indirizzo che interessa alla

- **OGNI ISTRUZIONE FA ACCESSO ALLA MEMORIA** cache, un riferimento
es. **ADD R1, R2, R3** genera 1 riferimento (alla memoria istruzioni) indirizzo fornito da cache
LW R1, 0(R2) genera 2 riferimenti (1 alla memoria dati + 1 alla mem. istruzioni)

fetch + accesso ai dati

istruzione genera un indirizzo che interessa la cache (FETCH) REFERENCE

istruzione genera 2 indirizzi che interessano la cache (FETCH, LOAD)

- Indichiamo quindi con $\bar{R}_{INST}$ = Average References per Instruction

$$\bar{R}_{INST} = \frac{N_{REF}}{N_{CPU}}$$

rif. complessivi
Numero Medio di riferim. per istruzione
istruz. dinamiche eseguite

- $N_{REF,i}$ = numero di riferimenti generati dalla istruzione i-esima

• $N_{REF,i1} = 1$ (es. istruz. ADD)

• $N_{REF,i2} = 2$ (es. istruz. LW)

ogni istruzione genera RIFERIMENTI alla cache, cioè indirizzi

$$\bar{R}_{INST} = \frac{N_{REF}}{N_{CPU}} = \frac{1}{N_{CPU}} \sum_{i=1}^{N_{CPU}} N_{REF,i}$$

- Nell'esempio precedente:

• N_{REF} = numero totale di riferimenti = 3

• $N_{CPU} = 2$

→ $\bar{R}_{INST} = 1.5$

in media genero 1,5 riferimenti alla cache

Nota: la cache vede al suo ingresso RIFERIMENTI, non istruzioni !

Equazione delle prestazioni MODIFICATA

- Equazione delle prestazioni:

$$T_{CPU} = C_{CPU} \cdot T_C = N_{CPU} \cdot \overline{CPI} \cdot T_C$$

N di istruz. dinamic. eseguite
Tempo clock
nel caso ideale, LRU e SW un solo ciclo di access
medio di cicli medi per istruzione

- $1/CPI$ stavolta è un "CPI-efficace" = "CPI ideale" +
 + numero cicli medi di stallo (penalty) per istruzione

cache non cattura tutto il working set

$$T_{CPU} = N_{CPU} \cdot \left(\overline{CPI}_{ideal} + \frac{\overline{R}_{INST}}{N_{CPU}} \cdot m \cdot C_{pen} \right) \cdot T_C$$

C PEN
Formib. dall'eq. ideale
riferim. di istruzione su cache
miss
penalty

$$= T_{CPU,ideal} + N_{REF} \cdot m \cdot C_{pen} \cdot T_C$$

Tempo aggiuntivo dovuto alla Penalty

→ Ho due modi per aumentare le prestazioni:

- Diminuire il miss rate
- Diminuire i cicli di penalty

$$\text{Tempo aggiuntivo per accedere a livelli successivi} = \frac{N_{rif} \cdot \text{Prob. miss}}{N_{min} \cdot \text{Tempo Penalty}}$$

Tempo in più che mi serve

Eg. mostos. ciclice che e' bene avere un clock più alto e eseguire un numero minore possibile di cicli per istruzione ma nel caso REALE e' anche bene DIMINUIRE MISS RATE e Cicli di Penalty

Esempio: impatto della cache sulle prestazioni

- Supponiamo di avere un processore a 200 MHz (5 ns per ciclo)

$\Rightarrow T_c = 5 \text{ ns}$

- $CPI_{ideal} = 1.1$

- 50% arit/log, 30% ld/st, 20% control Mix di istruzioni

- Supponiamo che il miss-rate per le istruzioni sia l'1% 50 cicli penalty

- Supponiamo che il 10% delle operazioni in memoria subisca miss con una penalty pari a 50 cicli (di clock) per l'accesso ai dati

- Da (**) si ha per il CPI: $CPI = CPI_{ideal} + (m_I + N_{CPU,L/S}/N_{CPU} \times m_D) \times C_{pen}$

$$= 1.1 \text{ (cycles)} + \overbrace{0.01 \text{ (miss/inst)} \times 50 \text{ (cycles/miss)}}^{C_{pen} \cdot m_{inst}}$$

$$+ \overbrace{0.30 \text{ (dmop/inst)} \times 0.10 \text{ (miss/dmop)} \times 50 \text{ (cycles/miss)}}^{C_{pen} = \text{miss}_{L/S} = 0.3 \text{ istru.}}$$

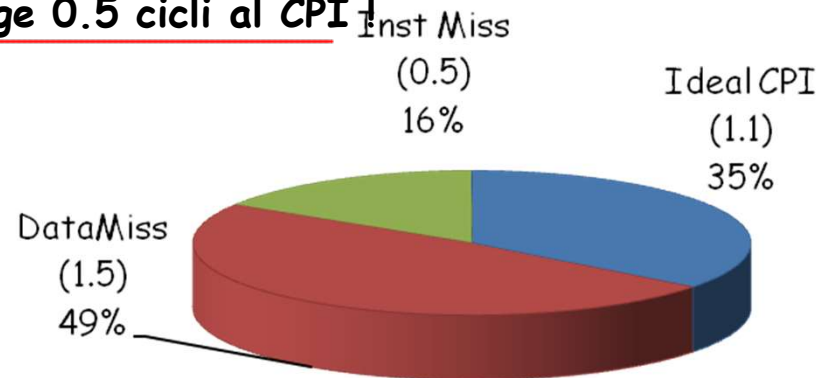
$$= 1.1 \text{ cycles} + 0.5 \text{ cycles} + 1.5 \text{ cycle} = \mathbf{3.1} \text{ CPI equivalente}$$

Fetch instr. + 1,5 cicli per accesso ai dati

→ Per il 65% (2.0/3.1) del tempo il processore attende la mem.!

→ Un 1% di miss-rate per le istruzioni aggiunge 0.5 cicli al CPI!

dmop=data memory operation



Con CACHE A ACCESSO DIRETTO, solo un blocco su 4 viene accettato (per es.) come può migliorare?

Le più indirizzi sono legati allo stesso slot

Patterson: 5.4

Cache Associative

ASSOCIO ALLO STESSO SLOT PIÙ BLOCCHI

- Sullo stesso indirizzo di cache insistono più blocchi. ASSOCIO ad uno stesso slot
→ posso fare in modo di «catturarne» più di uno

Il numero di blocchi associati ad un dato indirizzo di cache è fisso per una data cache ed è chiamato: dalla cache più blocchi

A = "grado di Associatività" della cache o "numero di vie" della cache

A = #BLOCCHI ASSOCIATI a quello SLOT ← numero quadrato

- Detti inoltre

- B = dimensione del Blocco di cache in byte
 - C = Capacità della cache in byte
-] Potenze di 2

dove:

$$A \in \{1, 2, \dots, A_{MAX}\}$$

A = 2, 2 blocchi in quello slot

$$\left\{ \begin{array}{l} B = 2^b \text{ con } b \in \{0, 1, 2, \dots, b_{MAX}\} \\ C = 2^c \text{ con } c \in \{0, 1, 2, \dots, c_{MAX}\} \end{array} \right\} \text{ multipli di 2}$$

SOTTOSISTEMI MEMORIA
dove CPU RAGIONANO a BLOCCHI

C capacità della cache

- Sussistono le seguenti relazioni

$$N_C = C / B = 2^{c-b} \rightarrow \text{Numero di blocchi della Cache}$$

$$N_M = M / B = 2^{m-b} \rightarrow \text{Numero di blocchi della Memoria}$$

(essendo M la dimensione della memoria in byte e $m = \log_2(M)$)

dim. memoria

Cache ragiona a Blocchi

CPU può fare local e STORE

Blocchi Cache a 32 o 64 byte

Cache Set Associative

Blocco del disco, es. 4 KB

Cache ad 1 via 1 via
(direct mapped o direct access)

accesso diretto

set	TAG	DATA
0		
1		
2		
3		
4		
5		
6		
7		

8 blocchi, cioè 8 set

- Data una cache ad 1 via, con 8 blocchi, posso organizzare gli stessi blocchi in modo diverso

Raggruppo i blocchi in insiemi o "set":
il numero di blocchi in uno stesso set è detto A="numero di vie" o associatività della cache

set prima era C/B

$$N_s = C / (B \cdot A)$$

set da indovinare dall'indirizzo dato dalla CPU

diviso per blocchi e per associatività

→ Numero di Set della cache

Aumento A, ad un set corrispondono A blocchi, # set = C/B/A

→ L'indirizzo del set sarà

l'indirizzo di X_M

$$X_s = X_M \bmod N_s$$

...e il tag

Preleva set di interesse

$$X_s = X_M \% \# \text{set}$$

l'indirizzo del blocco

$$X_T = X_M \div N_s$$

diviso per N_s per il Tag

trova vero Tag

Trova vie andate, nuovo set altro

Cache Set-Associativa a 2 vie (A=2)

set	TAG	DATA	TAG	DATA
0				
1				
2				
3				

a parità di memoria si riducono i set ma ho più possibilità nel mapping da memoria a cache

Cache Set-Associativa a 4 vie (A=4)

set	TAG	DATA	TAG	DATA	TAG	DATA	TAG	DATA
0								
1								

4 blocchi, 2 soli set

ad ogni set sono associati 4 slot

Nota: come caso particolare queste relazioni valgono anche per le cache ad accesso diretto (A=1 → N_s=N_c)

Cache Full Associative

Cache Set-Associativa a 8 vie (A=8) → Completamente Associativa (o "Full Associative")



col tag riconosco stato da memoria

$X_S = 1$ confronto X con Tag, tutti i bit sono di tag, tutti i blocchi vanno in set

$N_S = 1 \rightarrow X_T = X_M$ e $X_S = 0$ (cioè non occorre indirizzare il set) ho un solo set

confronto tutto eliminiamo con Tag

• Come individuo il blocco all'interno dello stesso set?

- Il problema si ha in generale in tutte le cache associative
- Devo confrontare in parallelo il tag X_T coi tag memorizzati T_k e loro validità
- La funzione di hit sarà più complessa

- $\forall k=1 \dots A, X_T == T_k \ \&\& \ V_k == 1 ? 1 : 0$

confronto
elimino
bit X_T specificato dal programma
Tag memorizzato nella via k-esima
è bit validità
devo fare un parallel
dei bit

Devo controllare il Tag $\forall A$

ora ogni set ha più blocchi, quindi più Tag

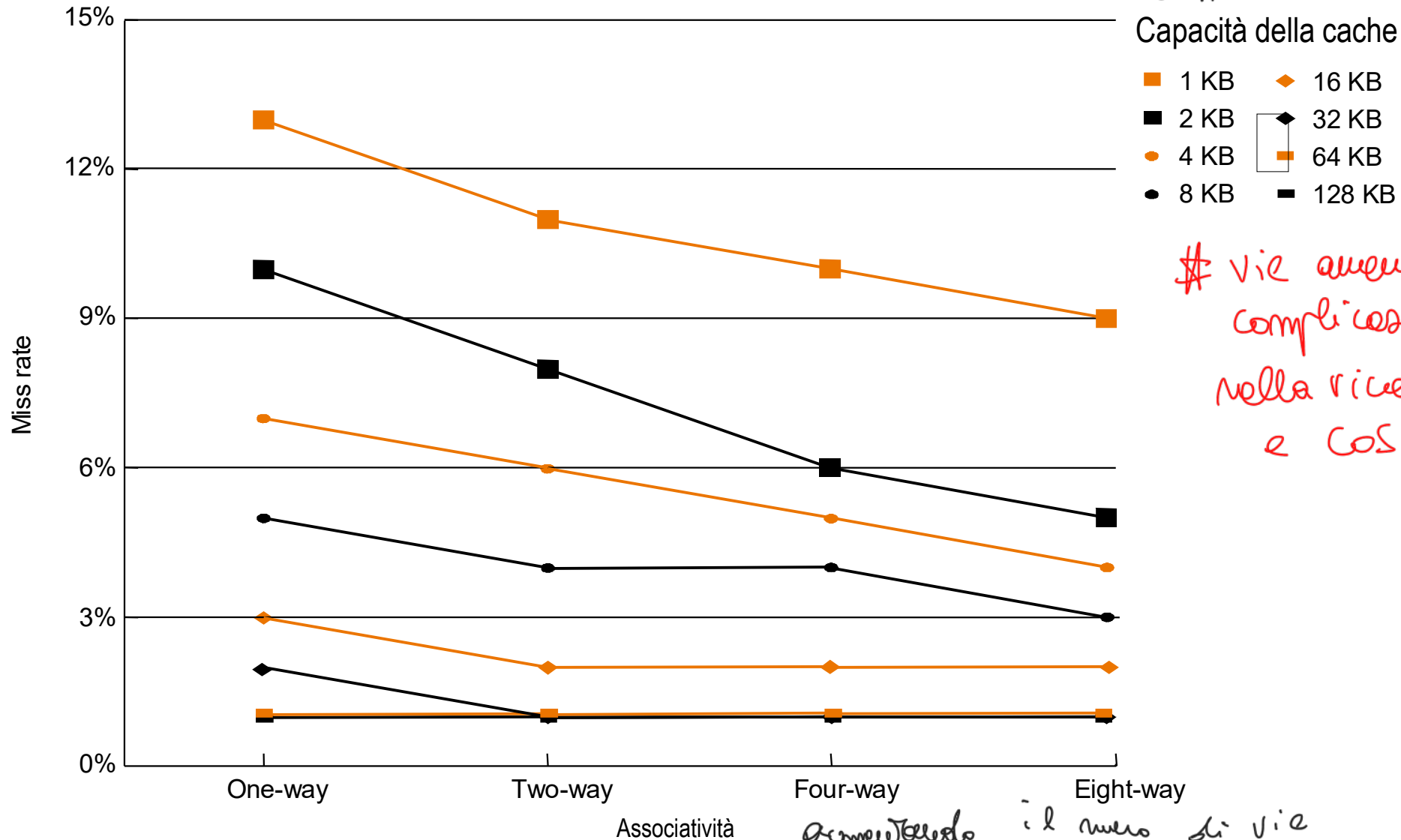
Valore la pena fare cache associative? SI!

Prestazioni: Direct-access vs. Set-associative cache

• E' utile ?

Sì: l'associatività tipicamente riduce il miss rate

*Aumentando Cache
diminuisce Miss
RATE
Capacità della cache*



*# vie aumenta
complicazione
nella ricerca
e COSTO*

*Aumentando il numero di vie
riduce il miss-rate*

Cerco un compromesso

Più blocchi sullo stesso set
Se devo scrivere su un blocco
già scritto, il blocco
di memoria viene aggiornato
devo copiare il vecchio
blocco in memoria

Politica di Rimpiazzamento (bits U)

• Come si procede quando un set è pieno?

devo salvare un blocco
ma set
è pieno

Aggiungo i vecchi blocchi al set

1) Scelgo il blocco da rimpiazzare secondo una certa politica

2) Carico il nuovo blocco nella locazione liberata

3) Leggo o scrivo in quel blocco

Se ne ho più di blocchi, come decido?

ELIMINO BLOCCO = SCRIVO IL SUO VALORE IN MEMORIA

• Le politiche di rimpiazzamento possono essere di svariato tipo;
noi consideriamo le seguenti tre:

elimino blocco = lo copio in
memoria

• **random**: scelgo un blocco a caso

- Veloce, economica, prestazioni scarse...

Tracce della LOCALITÀ TEMPORALE
LRU

• **LRU** (Least Recently Used): scelgo il blocco usato meno recentemente

Tempo dopo del
blocco usato più
recentemente

- Ottime prestazioni (sfrutta la località temporale), costosa da implementare
(con due vie è facile tenere traccia, ma per più vie l'hardware è più complesso)

- Spesso viene usato una versione semplificata

(es. Intel usa pseudo-LRU) che ha prestazioni un poco inferiori

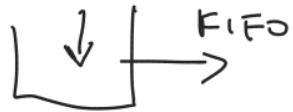
ALBERO DI PREONITÀ

• **FIFO**: scelgo il blocco caricato meno recentemente

primo
blocco che
intra, scarto

- Non molto usata

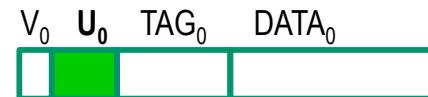
- Ha prestazioni intermedie fra random e LRU



Se FIFO ha meno
di primo
spesso,
non è
molto com.

• Per gestire la politica di rimpiazzamento si aggiungono ulteriori
bits (bits U) a fianco di ogni set

devo scegliere tra A vie, forse superiore
di $\log_2 A$



es. Per LRU occorrono $\lceil \log_2(A) \rceil$ bits U

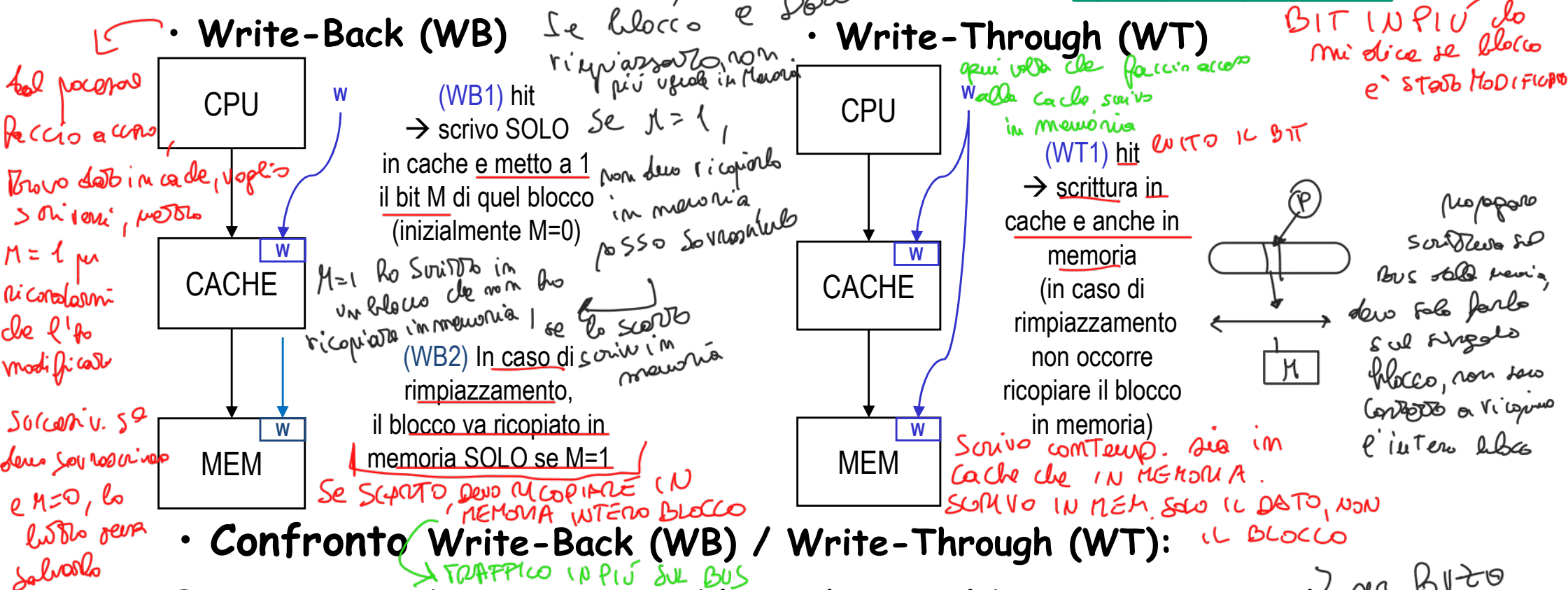
Politica di gestione delle scritture su hit (bit M)

Se ho un dato modificato? Voglio scrivere in cache, non in Memoria, su HIT devo risolvere

Devo scrivere in memoria in blocco che ho già in cache

- Un'ottimizzazione della cache consiste nell'evitare di ricopiare in memoria un blocco rimpiazzato, se NON è stato modificato (*)

- Ho due tecniche:



- WB: occorre un bit M per ogni blocco (in WT il bit M non occorre)
- + WB: riduco il traffico sul bus di memoria (write), rispetto a WT
- WB: si introduce un po' di traffico di aggiornamento della memoria

Schema architetturale di una cache 4-way set associative con politica di rimpiazzamento LRU/FIFO e write-back

- $M = 4 \text{ Gbyte}$ *32 bit per indirizz. Memoria*
- $A = 4$
- $B = 64 \text{ byte}$ *dim del blocco*
- $C = 32 \text{ Kbyte}$ *dim cache*

X indirizzo al byte

$X_M = X / B$
 $= X / 64$ *di resto per altro blocco*

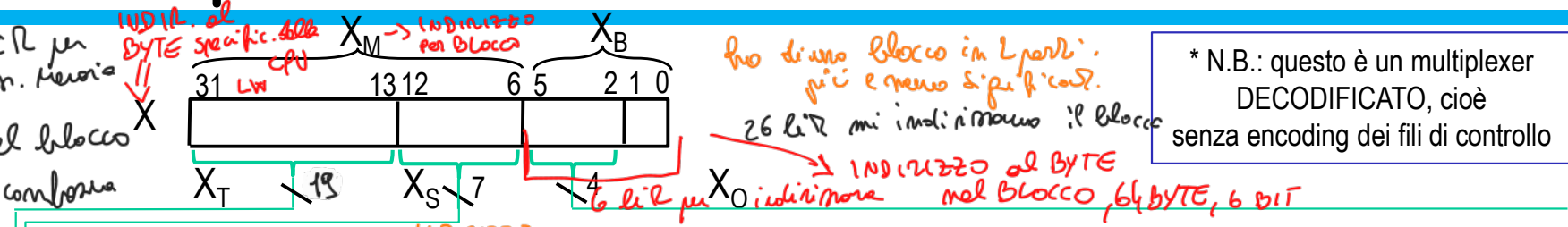
$N_S = C / B / A$ *#set*
 $= 2^{15} / 2^6 / 2^2$
 $= 2^7 = \text{ho } 128 \text{ righe}$

$X_S = X_M \bmod N_S$ *signific. di resto*
 $= X_M \bmod (C / B / A)$
 $= X_M \bmod 2^7$ *7 bit meno signific.*

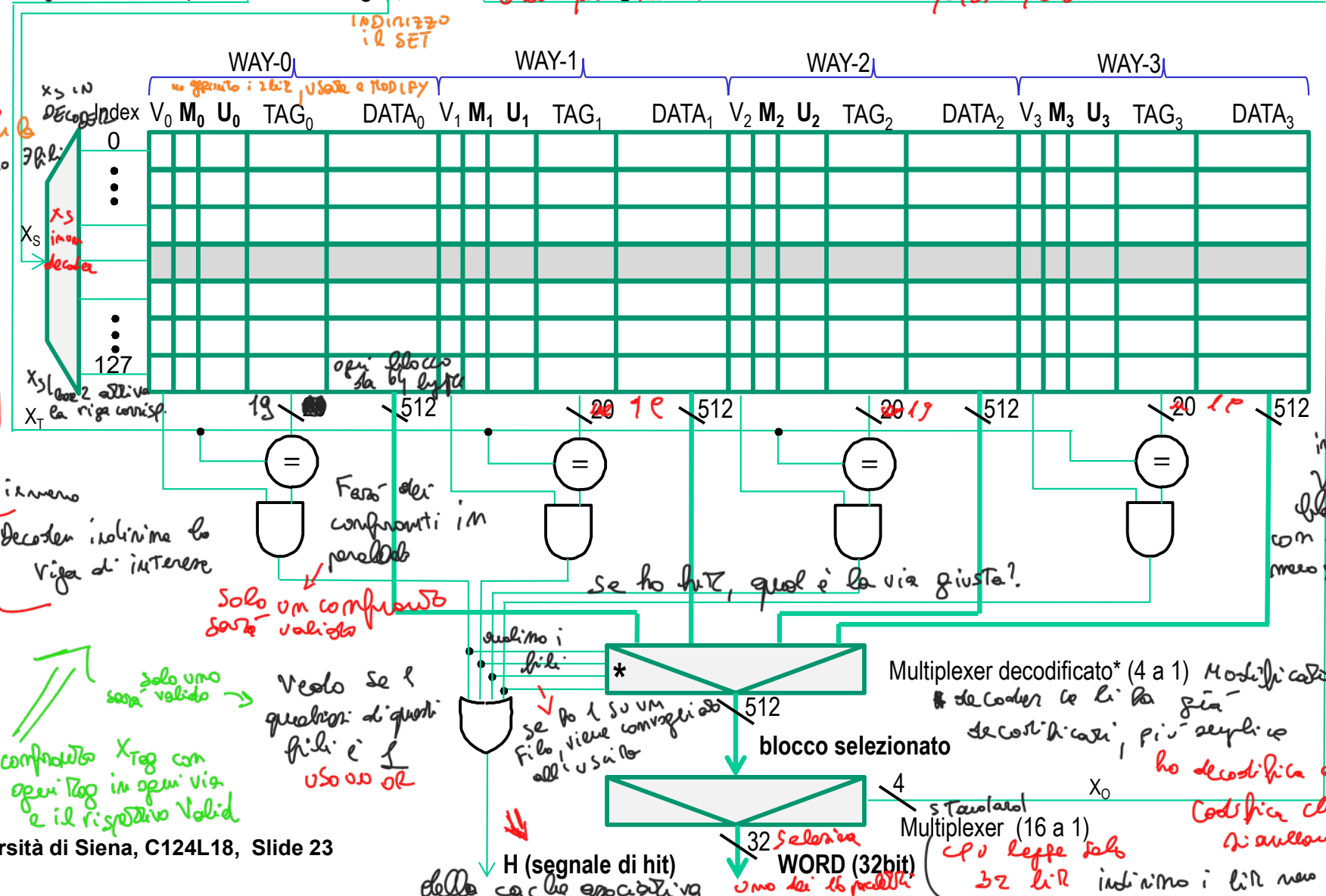
$X_T = X_M / N_S$
 $= X_M / (C / B / A)$
 $= X_M / 2^7$

X_0 = indirizzo word nel blocco

FUNZIONE DI CONFRONTO IN PARALLELO



* N.B.: questo è un multiplexer DECODIFICATO, cioè senza encoding dei fili di controllo



comparato X_{tag} con ogni tag in ogni via e il rispettivo valid

Vedo se i quali di questi fili è l'uso o no

Se ho 1 su un filo viene convalidato all'uscita

Se ho hit, qual è la via giusta?

Multiplexer decodificato* (4 a 1) Modificato decoder le linee già decodificate, più semplice ho decodifica e codifica che si annullano

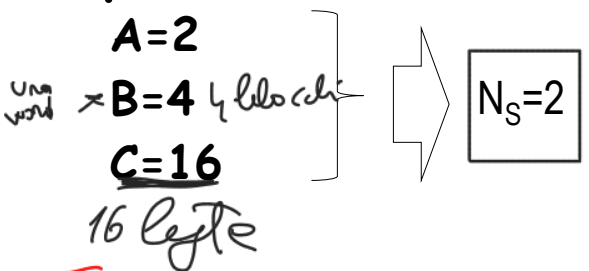
Multiplexer (16 a 1) *cpu legge solo 32 bit indirizzo i bit non sign.*

16 blocchi da 32 bit, ne esce 1 da 32 bit

seleziona una di 16 parole da 16 bit

Esercizio con politica di rimpiazzamento LRU (1)

- Consideriamo una cache set-associative a due vie, con una capacità totale di 4 word e blocchi da 1 word:



caratterizzare cache

	way0	way1
set 0		
set 1		

lasso con la
 parte dedicata
 all'indirizzo

- Supponiamo che i riferimenti X_M (al blocco) generati dal processore siano:

le due
sa due

R	0
W	6
R	0
R	1
W	4
R	0
R	2
R	3
W	5
W	4

es. colonna dell'indirizzo 0... 6...

trovare il V, D, B

seq. indirizzi

come vengono gestiti?

- Quanti hit e quanti miss saranno generati se la politica di rimpiazzamento è LRU?

come vengono gestiti i blocchi?

Esercizio con politica di rimpiazzamento LRU (2)

- Indirizzi 0, 6, 0, 1, 4, 0, ...

$$X_M \% 2^{N_s} / 2^{N_t} \text{ (divisione per 2)} \quad (N_s=2)$$

0: ($X_S=0, X_T=0$) → miss, carico nel set 0 (way 0) *Da cache è vuota* *sceglie arbitrariamente way 0* *1 via usata*

6: ($X_S=0, X_T=3$) → miss, carico nel set 0 (way 1) *non ho tag 3*
nessun caricamento, ma solo allungamento di un solo way 1 nel set 0

0: ($X_S=0, X_T=0$) → hit *ho dato in cache*
↳ ho usato way 0, LRU = way 2

1: ($X_S=1, X_T=0$) → miss, carico nel set 1 (way 0) *quasi su set 1*

4: ($X_S=0, X_T=2$) → miss, carico nel set 0 (way 1, **rimpiazzo il 3**) *ho inserito il tag 2* *sceglie LRU = way 1*

0: ($X_S=0, X_T=0$) → hit

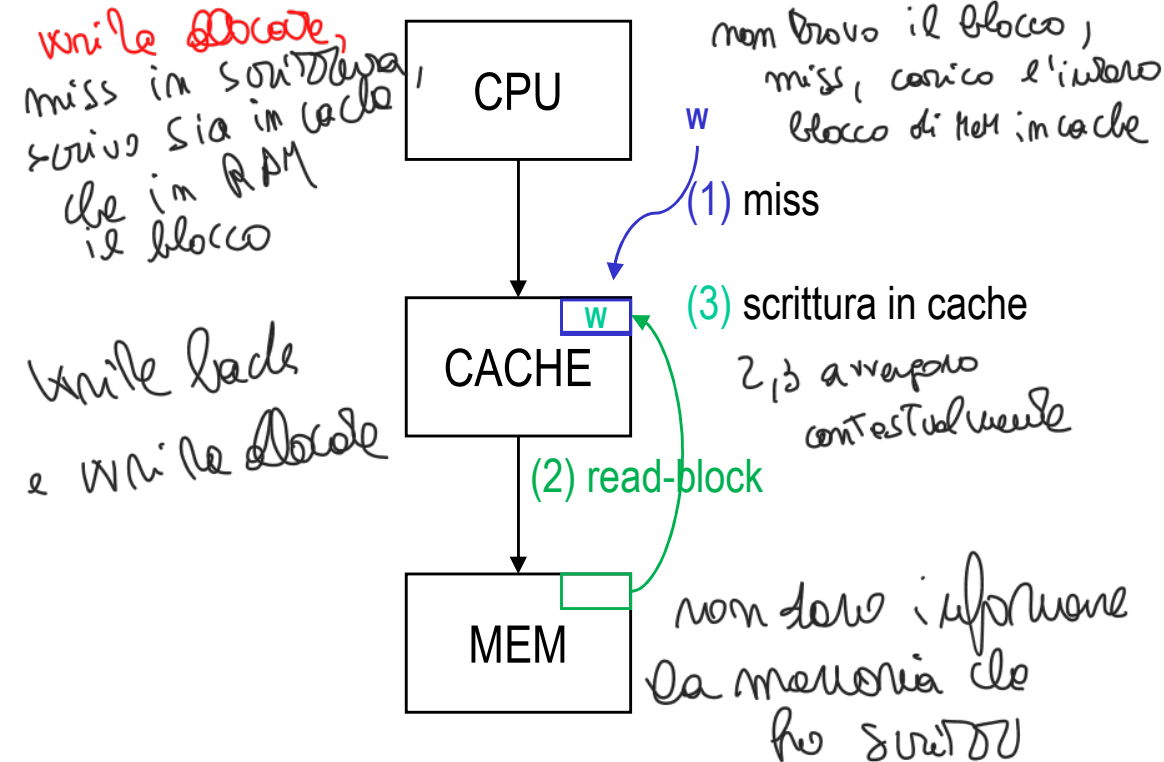
	way0	way1
set 0	0	iru
set 1		
set 0	0	3
set 1		
set 0	0	3
set 1		
set 0	0	3
set 1	0	iru
set 0	0	2
set 1	0	iru
set 0	0	2
set 1	0	iru

Politica di gestione delle scritture su miss: WA/WNA

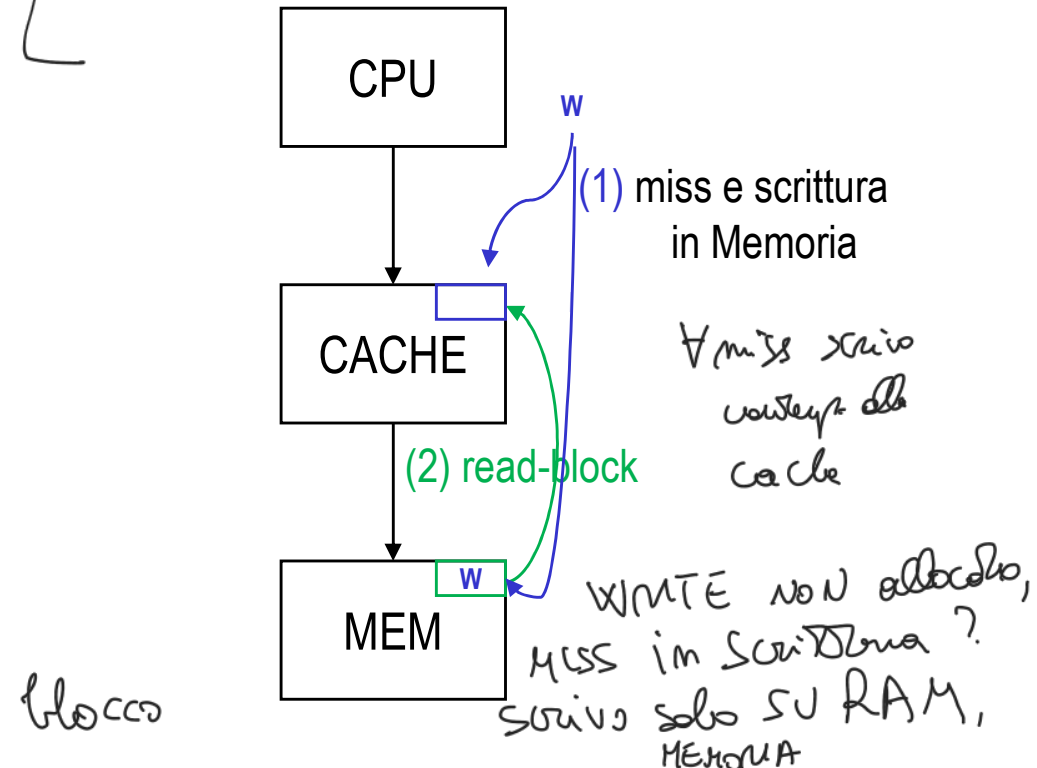
- Cosa faccio se durante una scrittura si verifica miss?

- Esistono due tecniche: *Allocate: tutto in cache un blocco prima di scrivere*

- **Write Allocate (WA)**



- tutto acceso miss WNA, logico prevede di inviare la scrittura alla cache e alla memoria*
- **Write Not-Allocate (WNA)**



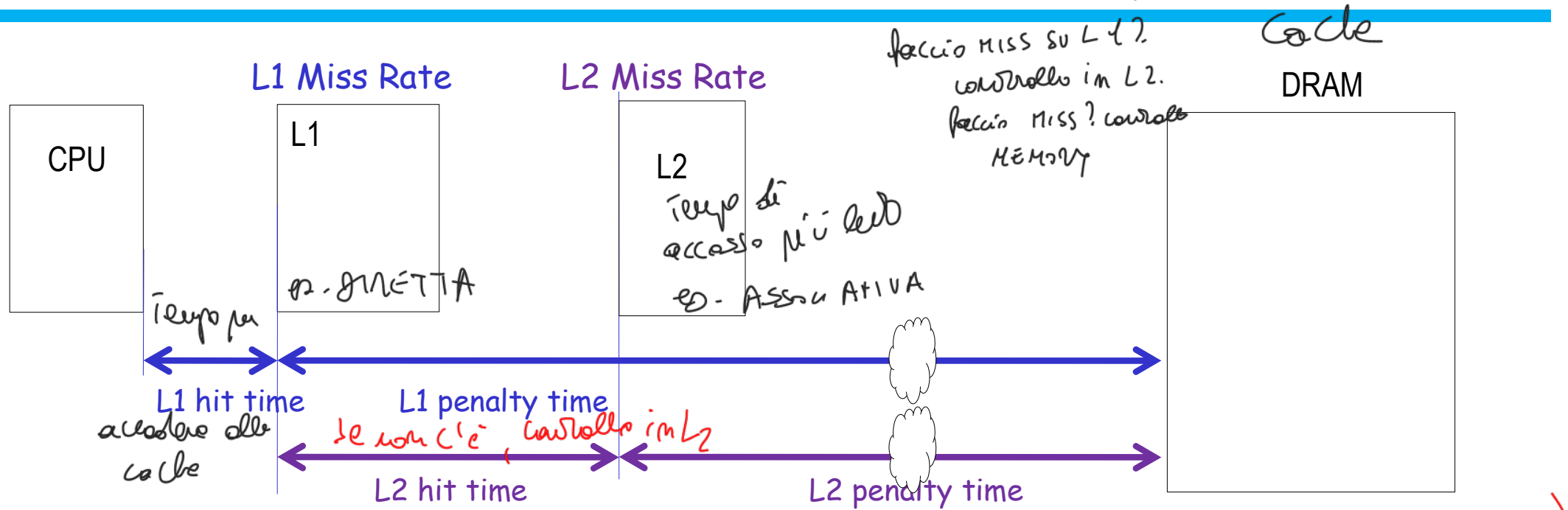
WA, NOTA: Le operazioni 2 e 3 vengono compiute come una singola azione (si scrive nel blocco durante il suo caricamento)

WNA, NOTA: L'operazione 2, per scelta progettuale, può non essere implementata (ovvero NON si carica il blocco in cache): il beneficio effettivo di prelevare il blocco o meno può dipendere fortemente dal tipo di applicazione; noi assumeremo che venga compiuta.

$$AMAT = t_{hit} + t_{pen} + m$$

Cache a 2 livelli

Struttura il sistema con più livelli di



$$AMAT = L1 \text{ Hit Time} + L1 \text{ Miss Rate} * L1 \text{ Penalty Time}$$

$$L1 \text{ Penalty Time} = L2 \text{ Hit Time} + L2 \text{ Miss Rate} * L2 \text{ Penalty Time}$$



$$AMAT = L1 \text{ Hit Time} + L1 \text{ Miss Rate} * (L2 \text{ Hit Time} + L2 \text{ Miss Rate} * L2 \text{ Penalty Time})$$

Tempo di accesso medio

Valori tipici

• L1 *es. di cache*

- Capacità: decine di KB
- Hit-time: uno (o due) cicli di clock
- Miss-rate tipici: 1-5%

più veloce, costoso

• L2: *es. associativa*

- Capacità: centinaia di KB
- Hit-time: qualche ciclo di clock
- Miss-rate tipici: 10-20%

veloce a sovrivere i miss della 1ª cache

peggioro perché nel primo livello c'è la località del program, nel secondo livello sono gli accessi "casuali" che non riescono a servire all'L1

• Località Frastagliata

• I miss di L2 sono una frazione dei miss di L1

- Perché allora il miss rate di L2 è così alto?

Sono più vicini all'applicazione, meglio le località

Cache a più livelli: esempio (confronto con 1 solo livello)

- Supponiamo di avere L1 e L2 con i seguenti dati:

- L1 Hit-Time = 1 cc
- L1 Miss-rate = 5%
- L2 Hit-Time = 5 cc
- L2 Miss-rate = 15% (% dei miss di L1 che fa miss in L2)
- L2 Penalty-Time = 100 cicli

- Con 2 LIVELLI:

- L1 Penalty Time = L2 Hit Time + L2 Miss Rate * L2 Penalty Time = $5 + 0.15 \times 100 = 20$ cc
- AMAT = L1 Hit-Time + L1 Miss-Rate * L1 Penalty-Time = $1 + 0.05 \times 20 = 2$ cc

20 cicli di Penalty Time

media mente

- Con 1 solo LIVELLO:

- L1 Hit-Time = 1 cc
- L1 Miss-rate = 5%
- L1 Penalty-Time = 100 cc
- AMAT = $1 + 0.05 \times 100 = 6$ cc

100 secondi

over più livello coperto
budget

→ La cache L2 migliora il tempo medio di accesso di un fattore 3

Formule, non pretendo che le ricordiamo

Contributi separati delle istruzioni generiche e L/S

$$T_{\text{CPU}} = N_{\text{CPU}} \cdot \left(\overline{CPI}_{\text{ideal}} + \left(\frac{N_{\text{REF},I}}{N_{\text{CPU}}} \cdot m_I + \frac{N_{\text{REF},D}}{N_{\text{CPU}}} \cdot m_D \right) \cdot C_{\text{pen}} \right) \cdot T_C$$

$$m_I = N_{\text{MISS_ISTRUZIONI}} / N_{\text{REF},I}$$

$$m_D = N_{\text{MISS_DATI}} / N_{\text{REF},D}$$

$$(**) = N_{\text{CPU}} \cdot \left(\overline{CPI}_{\text{ideal}} + \left(m_I + \frac{N_{\text{CPU,L/S}}}{N_{\text{CPU}}} \cdot m_D \right) \cdot C_{\text{pen}} \right) \cdot T_C$$

$$N_{\text{REF},I} = N_{\text{CPU}}$$

$$N_{\text{REF},D} = N_{\text{CPU,L/S}}$$

$$= T_{\text{CPU,ideal}} + (N_{\text{CPU}} \cdot m_I + N_{\text{CPU,L/S}} \cdot m_D) \cdot C_{\text{pen}} \cdot T_C$$

m = miss combinato dati+istruzioni

$$= T_{\text{CPU,ideal}} + N_{\text{REF}} \cdot \left(\frac{N_{\text{CPU}} \cdot m_I + N_{\text{CPU,L/S}} \cdot m_D}{\underbrace{N_{\text{REF}}}_{=m}} \right) \cdot C_{\text{pen}} \cdot T_C$$

$$m = \frac{N_{\text{CPU}} \cdot m_I + N_{\text{CPU,L/S}} \cdot m_D}{N_{\text{REF}}}$$

$$= T_{\text{CPU,ideal}} + N_{\text{REF}} \cdot m \cdot C_{\text{pen}} \cdot T_C = T_{\text{CPU,ideal}} + N_{\text{CPU}} \cdot m \cdot \overline{R}_{\text{INST}} \cdot C_{\text{pen}} \cdot T_C$$

Nota: per una split cache gli accessi avvengono in parallelo. Quindi: $m_{\text{I,SPLIT}} \neq m_I$, $m_{\text{D,SPLIT}} \neq m_D$, $m_{\text{SPLIT}} \neq m$

Fine parte teorica