

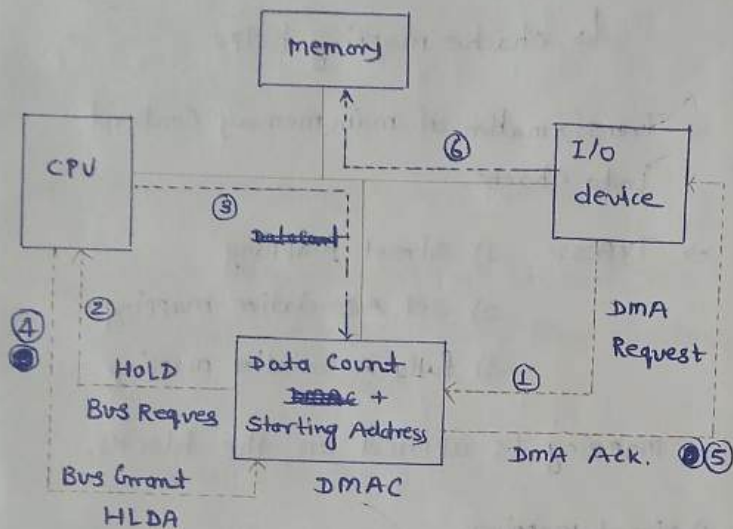
* Modes of Transfer :-

→ B/w I/O and CPU

• DMA (Direct memory Access) :-

→ Enables data transfer b/w I/O and memory Without CPU Intervention.

→ Need a Hardware — DMAC
— (DMA Controller)



• **Starting Address:** Memory Address starting from where the data transfer should be transferred.

• **Data Count:** No. of Bytes or Words to be transferred.

If memory is byte Addressable —

$$\text{Data Count} = \text{No. of Bytes}$$

If memory is Word Addressable —

$$\text{Data Count} = \text{No. of Words}$$

• During DMA Transfer CPU can perform only those μ -operation which do not require system buses.

⇒ if Data Count register has 4 bit means max data count = $(1111)_2$
= 15

DMAC is a special purpose processor which generates address and control signal for data transfer between memory and I/O.

• Modes of DMA Transfer :-

↳ for how much duration & when the control of buses should be given to DMAC.

- 1) Burst mode → old
- 2) cycle stealing mode → New
- 3) Interleaving mode

1) Burst mode :-

→ A burst of data (1KB-2KB generally) is transferred in one time before CPU takes back the control of the buses.

2) Cycle stealing mode :-

→ Slow I/O device takes some time to prepare a byte or word. During this time CPU keeps the control of the buses. Whenever the word or byte is ready then CPU give the control of buses to DMAC for one memory cycle. During this cycle the prepared byte or word is transferred with memory.

3) Interleaving mode :-

→ Whenever CPU performs Internal operation and does not require buses then DMAC will be given the control of the buses.

→ Hence CPU will Never be blocked due to DMA

- Time required to prepare the data in I/O = t_x
- Time required to transfer the data to memory = t_y

$$\% \text{ of time CPU is blocked due to DMA} = \frac{t_y}{t_x + t_y} \times 100$$

↳ for burst mode

$$\% \text{ of time CPU is blocked due to DMA} = \frac{t_y}{t_x} \times 100$$

↳ for cycle stealing mode

→ t_x is dependent on device's Internal Speed.

→ t_y
 ↳ given directly in question
 ↳ dependent on memory

* Locality of Reference :-

- If CPU has requested one Address for memory Access, then that particular Address or nearby Addresses will be accessed soon.

• Types -

(1) Temporal (based on Time)

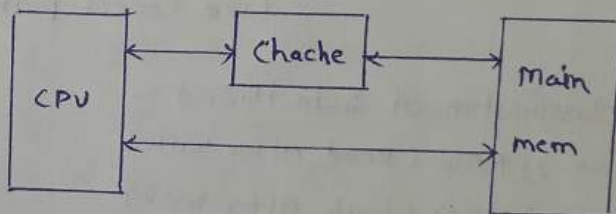
↳ If same address is referred soon.

(2) Spatial (based on Space)

↳ Nearby Address are referred soon.

* Cache memory :-

- Based on locality of reference Concept, current demanded localities are kept into a smaller & faster memory known as cache.
- use of cache reduces CPU memory Access time.



- cache Hit - CPU's demanded Content is present in cache.
- cache miss - CPU's demanded Content is Not present in cache.
- Performance of cache is given by Hit Ratio.

$$H = \frac{\text{No. of Hits}}{\text{Total mem. References}}$$

$$\text{Miss Ratio} = 1 - H$$

* Working of cache :-

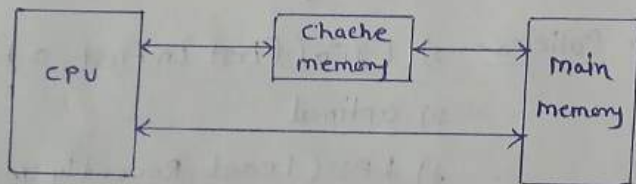
- first, CPU search the data in cache memory, if data found, then processes it.
- If data Not found in cache then search in primary memory and loads it into the cache.
- frequently Accessed data are always found.

* Average memory Access time -

$$T_{avg} = H * \text{Mem Access time for each Hit} + (1-H) * \text{mem Access time for each miss}$$

* Types of cache Access :-

1) Simultaneously Access (Parallel) :-



$$T_{avg} = H * t_{cache} + (1-H) t_{mem}$$

2) Hierarchical Access (serial) :- (Level wise)

$$T_{avg} = H * t_{cache} + (1-H) (t_{cm} + t_{mm})$$

$$= H * t_{cm} + t_{cm} - H t_{cm} + (1-H) t_{mm}$$

$$T_{avg} = t_{cm} + (1-H) t_{mm}$$

$$\text{Block Access time in main mem} = \text{Block size} \times t_{mm}$$

⇒ T_{avg} When block Transfer included -

1) Simultaneously -

$$T_{avg} = H * t_{cm} + (1-H) (t_{block} + t_{cm})$$

2) Hierarchical -

$$T_{avg} = H * t_{cm} + (1-H) (t_{cm} + t_{block} + t_{cm})$$

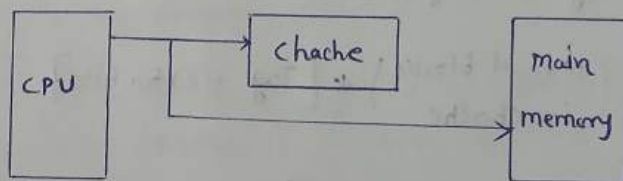
* Cache - Write :- (Write Propagation)

- If CPU update the Content in cache then the original content in ~~mem~~ main memory also should be updated.

→ Two methods -

- 1) Write Through
- 2) write back

1) Write Through :-



- The write operation is performed in cache and main memory simultaneously

→ ⊕ : 1) Data Consistency

→ ⊖ : 1) Time Consuming

↳ for a write opⁿ main memory is accessed irrespective of hit/miss

$$(T_{avg})_{Read} = H * t_{cm} + (1-H) t_{mm}$$

or

$$H * t_{cm} + (1-H) (t_{cm} + t_{mm})$$

$$(T_{avg})_{write} = \max[t_{cm} + t_{mm}]$$

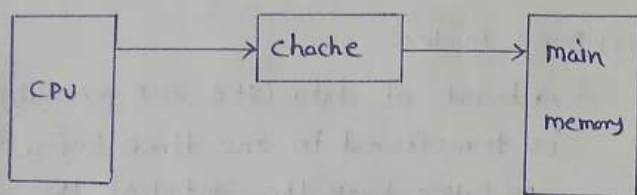
$$= t_{mm}$$

$$T_{avg} = \left(\frac{\text{fraction of Read}}{\text{of Read}}\right) * T_{Read} + \left(\frac{\text{fraction of Write}}{\text{of Write}}\right) * (T_{avg})_{write}$$

Effective Hit Ratio Rate -

fraction of Read * H for read

(2) Write Back :-



→ CPU updates only cache Content Whenever a block is replaced from cache then it is written back to main memory.

⊕ → 1) Time Saving

⊖ → 1) Data Inconsistency

$$(T_{avg})_{Read \text{ or } write} =$$

$$H * t_{cm} + (1-H) (t_{mm} + t_{write-back})$$

OR

$$H * t_{cm} + (1-H) (t_{cm} + t_{mm} + t_{write-back})$$

• write - Allocate:-

for a write miss, the missed block should be copied to cache from mem.

• write - No - Allocate:-

for a write miss, the missed block is not copied to cache.

→ for better performance -

write Through - uses No write Allocate

write Back - uses write Allocate

* cache mapping:-

→ Using main memory Address how cache memory is searched.

↳ cache mapping helps

→ Transformation of main memory Content into cache.

→ Types:-

- 1) Direct mapping
- 2) set Associative mapping
- 3) fully Associative mapping

→ mapping is applied on the blocks.

1) Direct mapping:-

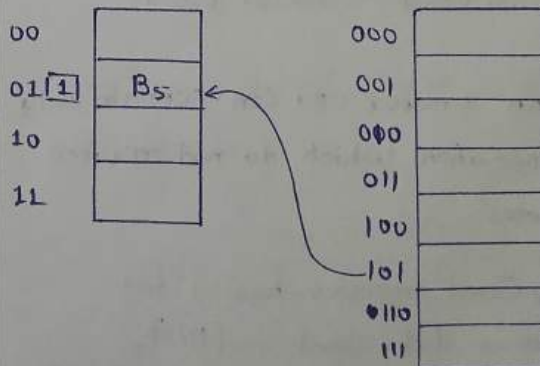
$$\frac{\text{No. of blocks/ Lines}}{\text{Line or Block size}} = \frac{\text{Cache size}}{\text{Line or Block size}}$$

• Physical Address is divided into 3 parts-

Tag	Number of cache Line	Byte offset
-----	----------------------	-------------

• cache Line Number = $\frac{M/m \text{ block No.}}{\text{No.}}$

$$\left(\frac{M/m \text{ block No.}}{\text{No.}}\right) \% \left(\frac{\text{Number of blocks in cache}}{\text{No.}}\right)$$



for $M/m = 101$

$$(101)_2 = 5, \quad 5 \% 4 = 1, \quad \begin{array}{cc} 1 & 0 & 1 \\ \downarrow & & \downarrow \\ \text{Tag} & & \text{Line No.} \end{array}$$

Tag	cm block No.	Byte offset
-----	--------------	-------------

$$\text{No. of bits in Byte offset} = \log_2(\text{block size})$$

$$\text{No. of bits in cm block No.} = \log_2 \left(\begin{array}{c} \text{No. of blocks} \\ \text{in cache} \end{array} \right)$$

$$\begin{aligned} * \text{ No. of tag Info stored in cache Controller} \\ = \text{No. of block in cache} \end{aligned}$$

$$* \text{ Total Tag size stored} =$$

$$\left(\begin{array}{c} \text{No. of blocks} \\ \text{in cache} \end{array} \right) * (\text{Tag-size})$$

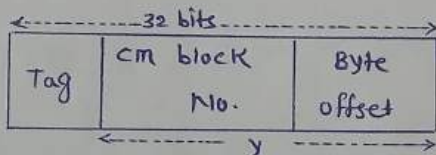
→ size of meta data

or

size of Tag directory

Q Consider a direct mapped cache of size 128 KB. The CPU generates of 32 bit. The Number of tag bits in main memory Address are ?

Ans -



$$y = \log_2(\text{cm size})$$

$$= \log_2(128 \text{ KB})$$

$$= 17 \text{ bits}$$

$$\text{Tag bits} = 15 \text{ bits}$$

$$* \text{ Index bits} = \text{No. of bits in cache mem block No.}$$

• Valid - Invalid bit

↳ There are one extra bit attach

In tag directory that verify that block fetches from main memory or Not?

$$\text{Valid bit} \begin{cases} \rightarrow 0 \rightarrow \text{Invalid bit} \\ \rightarrow 1 \rightarrow \text{valid bit} \end{cases}$$

• Write - back cache performance

↳ It can be improved if we can keep track of each cache block that which is modified & which one is not.

→ modified block only will require write back.

→ we use modified / dirty bit.

$$\text{dirty bit} \begin{cases} \rightarrow 0 - \text{only read} \\ \rightarrow 1 - \text{modified} \end{cases}$$

• Tag directory size =

$$\left(\begin{array}{c} \text{No. of blocks} \\ \text{in cache} \end{array} \right) * \left(\begin{array}{c} \text{Tag} \\ \text{bits} \end{array} + \begin{array}{c} \text{Extra} \\ \text{bits} \end{array} \right)$$

• T_{avg} for write-back cache-

$$(T_{\text{avg}})_{\text{read}} = (T_{\text{avg}})_{\text{write}}$$

$$= H * t_{\text{cm}} + (1-H) [t_{\text{block}} + \chi * t_{\text{block}}]$$

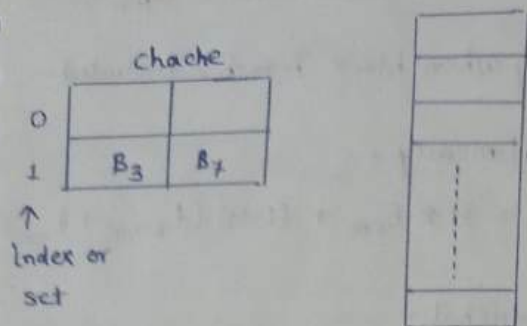
or

$$H * t_{\text{cm}} + (1-H) [t_{\text{cm}} + t_{\text{block}} + \chi * t_{\text{block}}]$$

χ : fraction of dirty

block, to be replaced from main memory

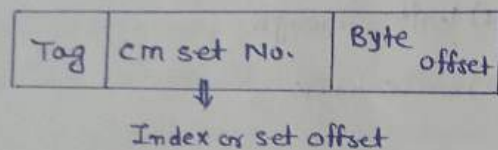
2) Set - Associative mapping :-



$$\begin{aligned} \text{* No. of Indexes} &= \frac{\text{No. of blocks in cache}}{\text{Associativity}} \\ \text{or} & \\ \text{No. of sets} & \end{aligned}$$

$$\text{* cm set No.} = \left(\frac{\text{mem block No.}}{\text{No.}} \right) \% \left(\begin{array}{l} \text{No of sets} \\ \text{in cache} \end{array} \right)$$

* M/m Address -



* Tag directory size -

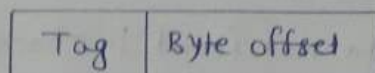
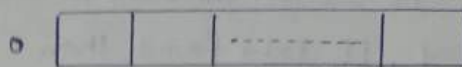
$$\left(\begin{array}{l} \text{No. of blocks} \\ \text{in cache} \end{array} \right) * [\text{Tag} + \text{Extra bits}]$$

* No. of bits in set offset =

$$\log_2 (\text{No. of sets in cache})$$

3) fully - Associative mapping :-

Set No. $\rightarrow 0$



Q If Cache Contains 2^{16} blocks and main memory contains 2^{24} blocks then size of cache tag directory size for fully set Associative cache is?

$$\begin{aligned} \text{Ans} \quad \text{m/m blocks} &= 2^{24} \\ &= 24 \text{ bits} \\ \text{tag bits} &= 24 \text{ bits} \end{aligned}$$

$$\text{Tag directory size} = 2^{16} * 24 \text{ bits}$$

* Block Replacement Policies :-

$\rightarrow$ Direct mapping does not require any replacement policy.

- $\rightarrow$ Policies -
- 1) FIFO (First in First out)
 - 2) Optimal
 - 3) LRU (Least Recently Used)

1) FIFO:- The block which Come first in cache should be replaced.

$\rightarrow$ Let 4 Way set Associative with ~~16~~ total 16 cache blocks -

request - (0, 255, 1, 4, 8, 133, 159, 216, 129

63, 8, 48, 32, 73, 92, 155)

$\downarrow$
H

S ₀	0 48	4 32	8 92	216
S ₁	1	133	129	73
S ₂				
S ₃	255 155	3	159	63

(2) Optimal Replacement :-

→ Replace a block which is not going to be used in future for very very long time.

Block Reques - 2, 3, 4, 7, 6, 3, 4, 7, 5, 4, 7, 8

0	2 6
1	3 5
2	4 8
3	7

↓ ↓ ↓ ↓
H H H H

(3) LRU (Least Recently used) :-

→ Replace a block which has not been referred since long time.

→ 4way set Associative with 16 cache blocks

↳ Request : (0, 255, 1, 4, 3, 8, 133, 159, 216, 129, 63, 8, 48, 32, 73, 92, 155)

0, 255, 1, 4, 3, 8, 133, 159, 216, 129, 63, 8, 48, 32, 73, 92, 155

92, 155

00	0 48	4 32	8	216 92
01	1	133	129	73
10				
11	255 155	3	159	63

* Cache miss Penalty :-

→ Time Required to bring a missed block from main memory to ~~the~~ Cache.

→ Let

Cycle Required to send Address to memory = 1 cycle

Cycle require to access 1 main memory cell = 5 cycle

Cycle Require to transfer 1 cell data to cache = 1 cycle

• if cache block size = 4B

main memory block size = 1B

$$\Rightarrow \text{miss penalty} = 1 + (4 \times 5) + (4 \times 1) = 25 \text{ cycles}$$

→ Types -

- 1) Cold or Compulsory miss
- 2) Capacity miss
- 3) Conflict miss

(1) Cold or Compulsory miss

↳ First time access of a block will always cause a miss.

↳ To Reduce cold miss -

Increase Block size

(2) Capacity miss :-

→ If it's not a Cold miss and miss occurs because cache is full.

→ To reduce Capacity misses :-

"Increase the cache size"

(3) Conflict miss :-

- ↳ It's not a Cold and Capacity miss but due to the set is full and miss occurs.
- ↳ To reduce Conflict misses -
- No Conflict miss in fully Associative mapping.

• multilevel cache :-

1) Simultaneous Access -

$$T_{avg} = H_1 * t_1 + (1-H_1) H_2 * t_2 + (1-H_1)(1-H_2) * t_{mm}$$

2) Hierarchical Access -

$$T_{avg} = H_1 * t_1 + (1-H_1) H_2 * (t_1 + t_2) + (1-H_1)(1-H_2) * (t_1 + t_2 + t_{mm})$$

or

$$= t_1 + (1-H_1) [t_2 + (1-H_2) t_{mm}]$$

