

"Non-Linear Pipeline" :-

* Non-Linear Pipeline contain forward and backward connections so Reservation table is used to process the inputs in the Pipeline.

* Let us consider the following Reservation table to process the inputs in a Non-Linear pipeline.

	1	2	3	4	5
S ₁	X				X
S ₂		X		X	
S ₃			X		

- * multiple check-marks in the Row indicates that, Repeated use of some stage in different cycles
- * Contiguous check-marks in the Row indicate, extended use of some stage upto some number of cycles.
- * multiple check-marks in the column indicate, parallel use of a different stages in the same cycle.

Latency Analysis :-

- * Latency means clock cycle difference between the two successive initiations in the Pipeline it is a positive constant.

* Linear pipeline Latency is always 1 but Non-linear pipeline Latency will be depends on the reservation table.

* In the Non-~~Linear~~ linear pipeline some of the latencies causes the collision called 'forbidden latency' (Non-premissible latency).

* Some of the Latencies does not causes collision named as Non-forbidden Latency (Premissible Latency).

* To Detect the forbidden latency we need to take the difference between any two check-marks in the Row.

eg:-
 Row 1 : $(5-1) = 4$ ← Latency { also forbidden latencies }
 Row 2 : $(4-2) = 2$ ← Latency
 Row 3 : 1

* Based on the forbidden latency, collision vector is created that is :

$$C_i : [c_n \dots c_3 c_2 c_1]$$

$c \rightarrow \begin{cases} 0 : \text{Premissable} \\ 1 : \text{Non Premissable} \end{cases}$

$$\text{Collision vector} : [c_5 c_4 c_3 c_2 c_1]$$

$$= [0 \ 1 \ 0 \ 1 \ 0]$$

* Latency sequence is generated based on the collision vector to find out the Latency cycle. i.e.

1. Latency sequence with respect to " C_1 "

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
S_1	1	2			1	2	3	4			3	4	5		
S_2		1	2	1	2			3	4	3	4			5	
S_3			1	2					3	4					5

Latency cycle : $(2-1) = 1$

$(7-5) = 5$

$(8-7) = 1$

$(13-8) = 5$

$(1, 5), (1, 5), (1, 5) \dots$

Avg latency = $\frac{(1+5)}{2} = 3$

2. Latency sequence with respect to " C_3 "

	1	2	3	4	5	6	7	8	9	10	11	12
S_1	1			2	1		3	2		4	3	
S_2		1		1	2		2	3		3	4	
S_3			1			2			3			4

Latency cycle :-

$(4-1) = 3$

$(7-4) = 3$

$(10-7) = 3$

$3, 3, 3, 3 \dots$

Avg latency = 3

③ Latency sequence w.r. to "C₅"

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
S ₁	1				1	2	3			2	3	4	5		
S ₂		1		1			2	3	2	3			4	5	4
S ₃			1					2	3					4	5

Latency cycle :-

$$(6-1) = 5$$

$$(7-6) = 1$$

$$(12-7) = 5$$

$$(13-12) = 1$$

$$(S_1, 1), (S_1, 1), (S_1, 1)$$

$$\text{Avg Latency} = \frac{5+1}{2} = \underline{3}$$

$$\text{MAL [min Average Latency]} = \underline{3}$$

"Memory organization" :-

Types :-

Based on the style of accessing the memory system, memory organization is two kinds :

- 1.) Simultaneous access m/y organization
- 2.) Hierarchical access m/y organization

- Simultaneous access m/y organization :-

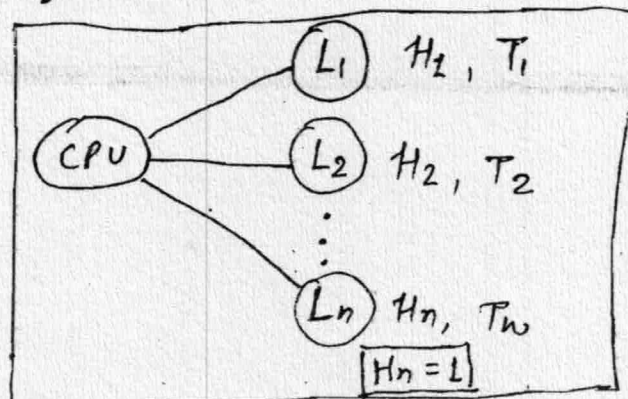
* In this organization, CPU directly connected to all the level of memories but accessing the memories in a sequence that is :

When there is a miss in level-1' memory then

* CPU directly access the data from the level-2' memory without copying into level-1'

* When there is a miss in level-2' m/y then CPU directly access the data from the level-3' memory without copying into level-2' and level-1' and so on.

* In this memory organization data block is need not be required in the level-1 memory



Here,

* $H_1, H_2, \dots, H_n \rightarrow$ Hit ratios

$H_1, H_2, \dots, H_n$ are the hit ratios of respective memories,

$$\text{Hit Ratio, } H = \frac{\# \text{ hits}}{\text{Total \# access}}$$

$$H = \{0 \text{ to } 1\}$$

Here

$T_1, T_2, \dots, T_n$ are the access times of a respective memories.

* Amount of the Time Required to access 1-word data from the memory is called as Average access time that is,

$$T_{\text{avg}} = H_1 T_1 + (1-H_1) H_2 T_2 + (1-H_1)(1-H_2) H_3 T_3 + \dots + (1-H_1)(1-H_2)\dots(1-H_{n-1}) H_n T_n$$

1-word takes — T_{avg}

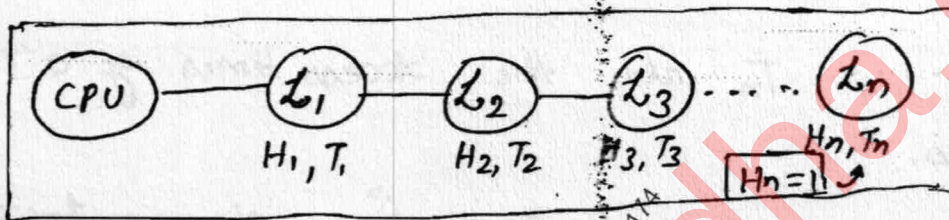
? # words — 1 sec.

$$\eta_{\text{memory}} = \frac{1}{T_{\text{avg}}} \text{ words/sec.}$$

- Hierarchical Access Memory organization :-

* In this organization, CPU always accessing the data only from the level-1 memory.

When there is a miss in level-1 mly then the Respective data-block is transferred from Higher-levels to lower level $\{L_i\}$. Later CPU accessing the data only from the level-1 memory.



$$T_{avg} = H_1 T_1 + (1-H_1) H_2 (T_2 + T_1) + (1-H_1)(1-H_2) H_3 (T_3 + T_2 + T_1) + \dots + (1-H_1)(1-H_2) \dots (1-H_{n-1}) H_n (T_n + T_{n-1} + \dots + T_1)$$

$$\eta_{mly} = \frac{1}{T_{avg}} \text{ words/sec.}$$

Code : I_1 : Data (L_2 mly)

I_2 :

I_3 : Data

I_4 :

I_5 : Data

I_6 :

I_7 : Data

} Reuse the data

Code:

Simultaneous

$$I_1: ML_1, T_2$$

$$I_3: ML_1, T_2$$

$$I_5: ML_1, T_2$$

$$I_7: ML_1, T_2$$

$$I_9: ML_1, T_2$$

5 times T_2 only
is referenced

Hierarchical

$$I_1: ML_1, T_2 + T_1$$

$$I_3: T_1$$

$$I_5: T_1$$

$$I_7: T_1$$

$$I_9: T_1$$

Reuse the
Data Space

} Locality of
Reference }

$$(1 * T_2) + (5 * T_1)$$

* In the Hierarchical memory design, Locality of reference is present so average accessing time is minimized.

Locality of Reference means CPU performing the operations on a Re-usable data-space.

It is two kinds :-

1) Temporal Locality : means same word in the Data-Block will be referenced by the CPU again and again in near future.

2) Spatial Locality : means adjacent words in the Data-block is referenced by the CPU in a sequence.

* To satisfy the above locality of references we need to design the Block-size with multiple words.

$$\boxed{\text{Block size} > \text{word size.}}$$

Note:- In the Computer system design m/y system designed using the Hierarchical principal to minimize the 'Speed gap' b/w CPU & m/y.

Note:- when the Question contain Hierarchy word, Hierarchy meaning or cache memory word then use Hierarchical memory organization otherwise use Simultaneous m/y organization.

Que:- Consider 2-Level m/y organization where Level-1 memory is 12 times faster than level-2 m/y and its access time is 30 ns less than the avg access time. Let Level 1 m/y access time is 40 ns what is the Hit ratio.

Level 1 = 12 level 2 m/y.

$$L_1 \rightarrow H_1, T_1$$

$$L_2 \rightarrow H_2, T_2$$

$$T_{avg} = H_1 T_1 + (1-H_1) H_2 T_2$$

$$(H_2 = 1)$$

$$70 = (H_1 \times 40) + (1-H_1) 1 \times 480$$

2 level m/y org.

$$70 = (H_1 \times 40) + (1-H_1) 480$$

$$H_1 = 0.93$$

$$d = \frac{T_2}{T_1}$$

$$12 = \frac{T_2}{T_1}$$

$$T_2 = 12 T_1$$

$$T_1 = T_{avg} - 30 \text{ ns}$$

$$T_{avg} = T_1 + 30 \text{ ns}$$

Let

$$T_1 = 40 \text{ ns then}$$

$$T_2 = 12 \times 40 \text{ ns} = 480$$

$$T_{avg} = 40 + 30 = 70 \text{ ns}$$

Que:- 3-level mly organization has the following specification:

Level	Access time/word	Block size in words	Hit ratio	Block Access time
1	20 ns	—	0.7	$T_1 = 20 \text{ ns}$
2	100 ns	2	0.9	$T_2 = 200 \text{ ns}$
3	200 ns	4	1	$T_3 = 800 \text{ ns}$

If the referred block is not in L_1 , then Transfer it from L_2 to L_1 . If Not in L_2 then L_3 to L_2 and L_2 to L_1 . How long will it take to access the Data Block.

:- Hierarchical: 3 level

$$\begin{aligned}
 T_{avg} &= H_1 T_1 + (1-H_1) H_2 (T_2 + T_1) + (1-H_1)(1-H_2) H_3 (T_3 + T_2 + T_1) \\
 &= 0.7 * 20 + (0.3) * 0.9 (220) + (0.3)(0.1) 1 (1020) \\
 T_{avg} &= 104 \text{ ns}
 \end{aligned}$$

* Que:- In a 2-level mly design Level-1 mly ~~access time~~ design Level-2 mly access time is 100 ns and its Bandwidth is 1 Gbps. Blocksize in the mly design is 128 Bytes. How much time is required to access the block from level-2 mly.

$$\frac{128}{2^{10}} = \frac{2^{20}}{2^{10}} \times 100 \text{ ns} = 2^{10} \times 100 \text{ ns} = 102400 \text{ ns}$$

$$\text{Total time} = \underset{\downarrow}{\text{Access time}} + \text{Transfer time} \left(\begin{array}{l} \text{If B.w.} \\ \text{given} \end{array} \right)$$

$$= 100 \text{ ns} + \frac{16 \text{ B}}{128 \text{ B}} \rightarrow 1 \text{ sec.}$$

$$\Rightarrow \frac{128 \text{ B}}{16 \text{ B}} \text{ sec.}$$

$$= 128 * 10^{-9} \text{ sec.}$$

$$= 100 \text{ ns} + 128 \text{ ns}$$

$$= 228 \text{ ns}$$

"Memory Hierarchy design" :-

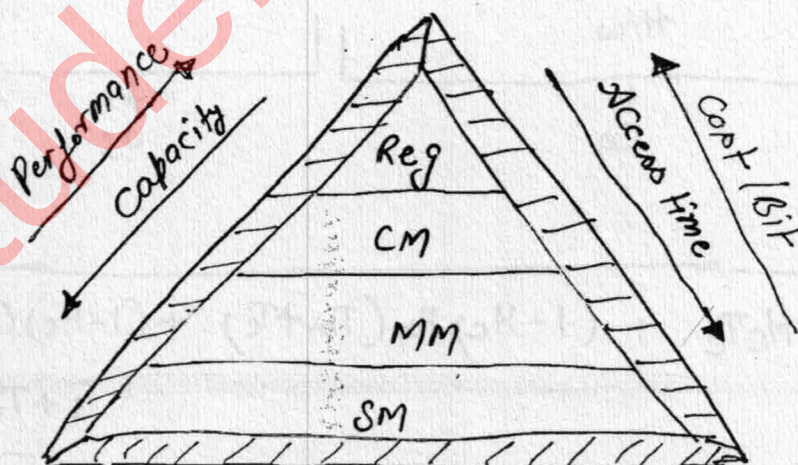
* In the Computer system design memory system is organized by using the memory hierarchy policy. According to a memory hierarchy design, system supported memory standards is as follows:

memory standards:

According to a M/y Hierarchy design accessing sequence of a system supported memories is as follows:-

Level :	1	2	3	4
Name	Register	Cache memory	main memory	Secondary memory
	(	(CM)	(mm)	(SM)
Typical Size :	< 1KB	< 16MB	< 16GB	> 100GB
Implementation :	customized multiports	SRAM (flip-flop)	DRAM (capacitors)	magnetic
Access time :	(0.25 - 0.5)	(5000 - 10000) (0.5 - 25)	(80 - 250)	500,000
Bandwidth (MB/sec) :	(20000 - 100000)	(5000 - 10000)	(1000 - 5000)	(20 - 150)
Managed by :	Compiler	HW	OS	OS
Backed by :	CM	MM	SM	Compact disk (CD)

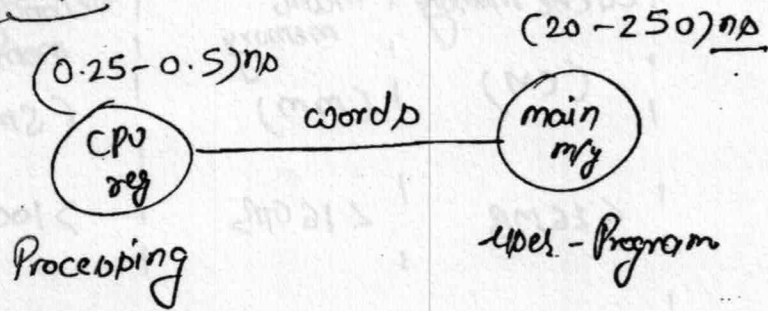
Bottom up approach:-



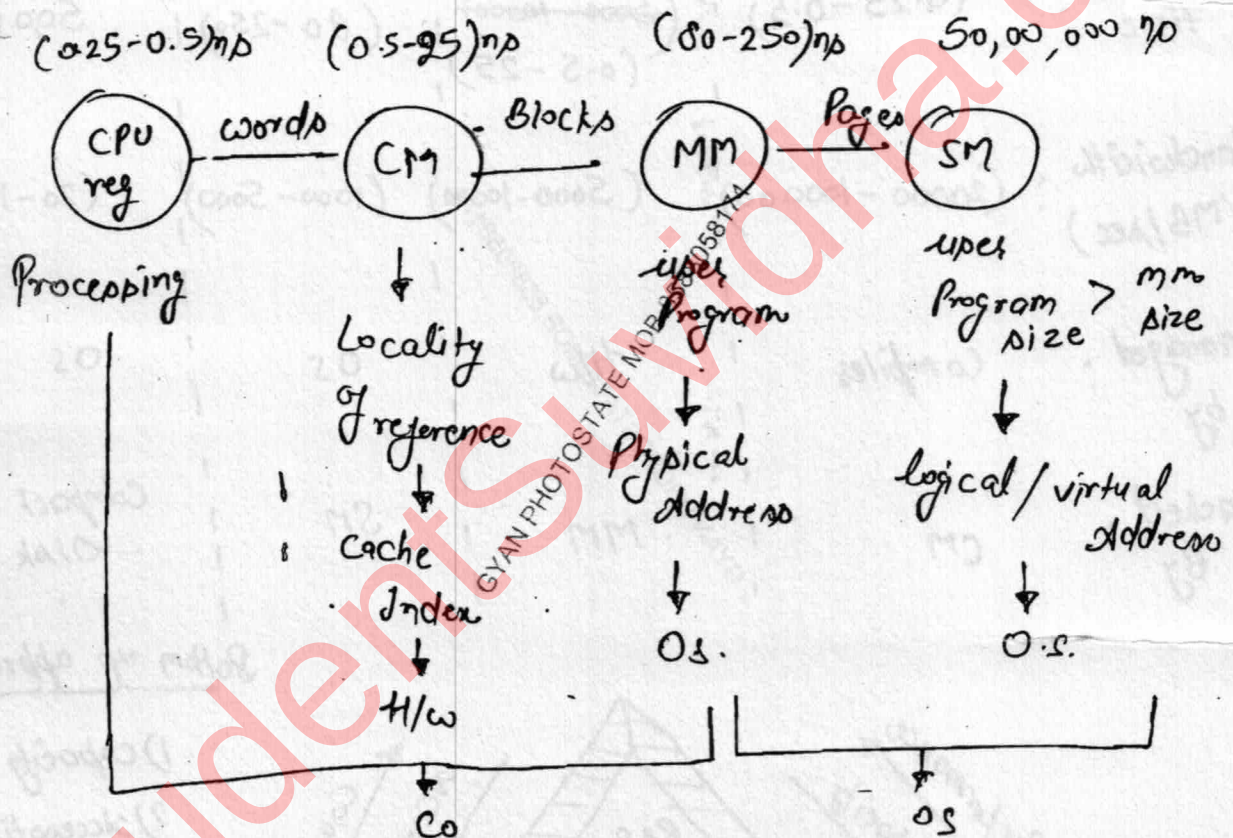
m/y Hierarchy Design

- 1) capacity ↓
- 2) access time ↓
- 3) Performance ↑
- 4) Cost/bit ↑
- 5) Expensive ↑

System old :



System new :



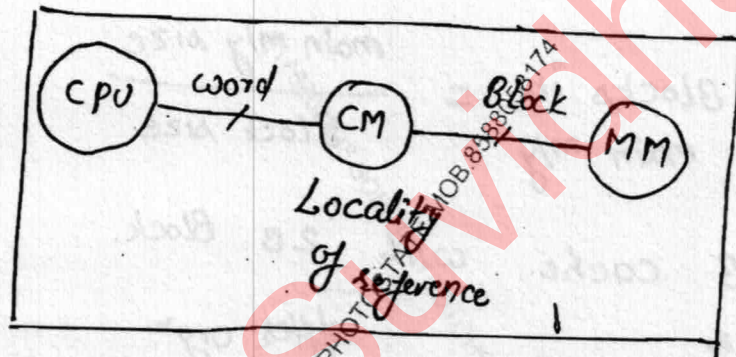
$$T_{avg} = H_c T_c + (1 - H_c) H_m (T_m + T_c) + (1 - H_c) (1 - H_m) H_s (T_s + T_m + T_c)$$

$[H_s = 1]$

3-level m/y structure

Cache Memory :-

:- Cache memory acts as an intermediate memory between CPU and main-memory, used to hold the image of a main-memory data. So CPU will be accessing the main-memory data from the cache-memory within less amount of time. Therefore speed-gap is synchronized between the CPU and main-memory.



Design Elements :-

- ① memory organization
- ② mapping techniques
- ③ Replacement algorithms
- ④ updating techniques
- ⑤ multilevel cache design.

memory organization :-

According to Hierarchy design, data will be transferred from main-memory to cache memory in a form of blocks as both of the memories are organized into blocks based on the block-size, as

$$\# \text{ Blocks (lines) in CM.} = \frac{\text{C.M. Size}}{\text{Block Size}}$$

$$\# \text{ Blocks in main m/y} = \frac{\text{main m/y size}}{\text{Block size}}$$

→ Consider 8B cache with 2B Block.

Before orgⁿ

8B → CM.

000	B
001	B
010	.
011	.
100	.
101	.
110	.
111	B

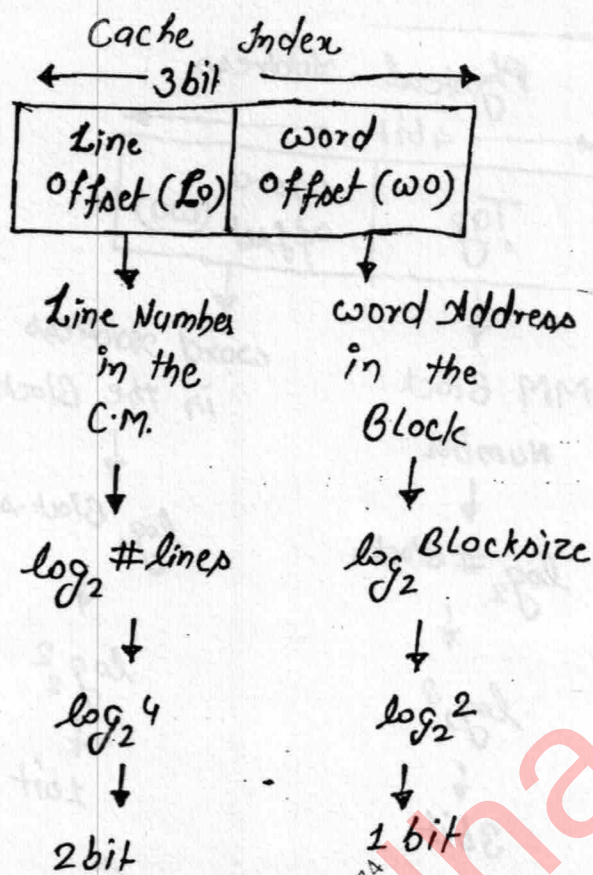
Cache index

After orgⁿ

$$\bullet \# \text{ lines (N)} = \frac{8}{2} = 4$$

$$\bullet 4 * 2B = 8B$$

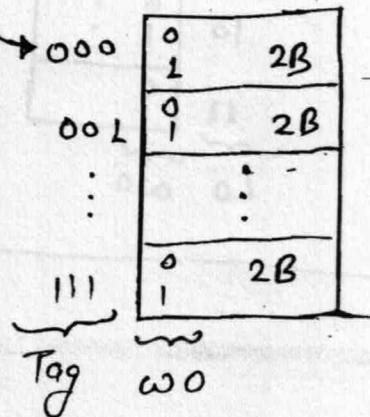
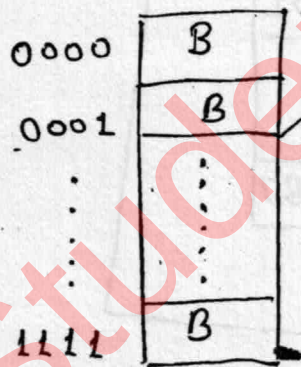
00	0	1	2B
01	0	1	2B
10	0	1	2B
11	0	1	2B
Lo	Wo		

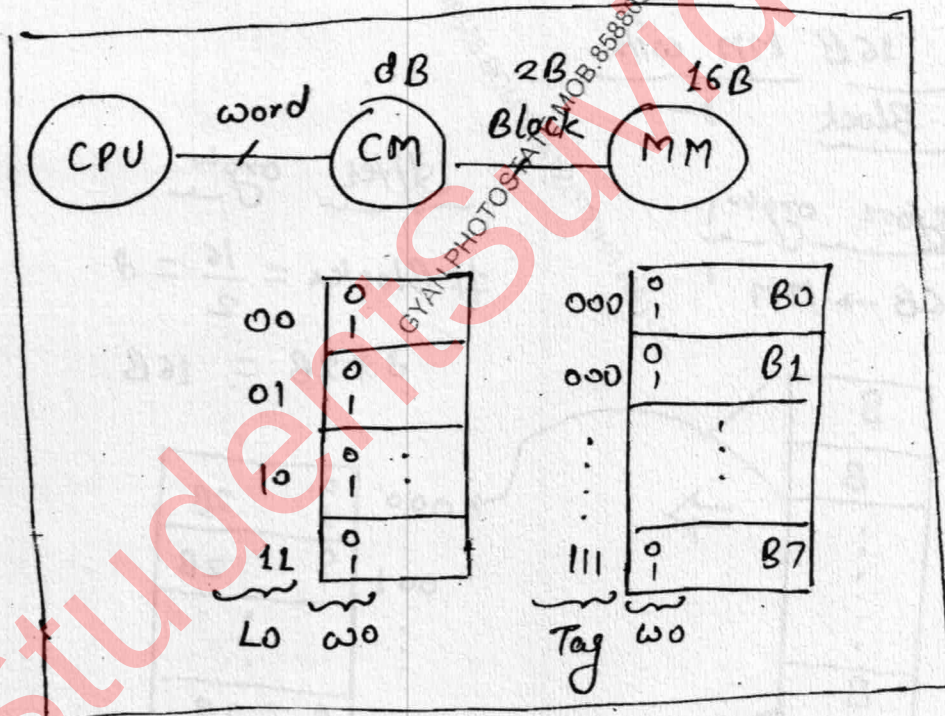
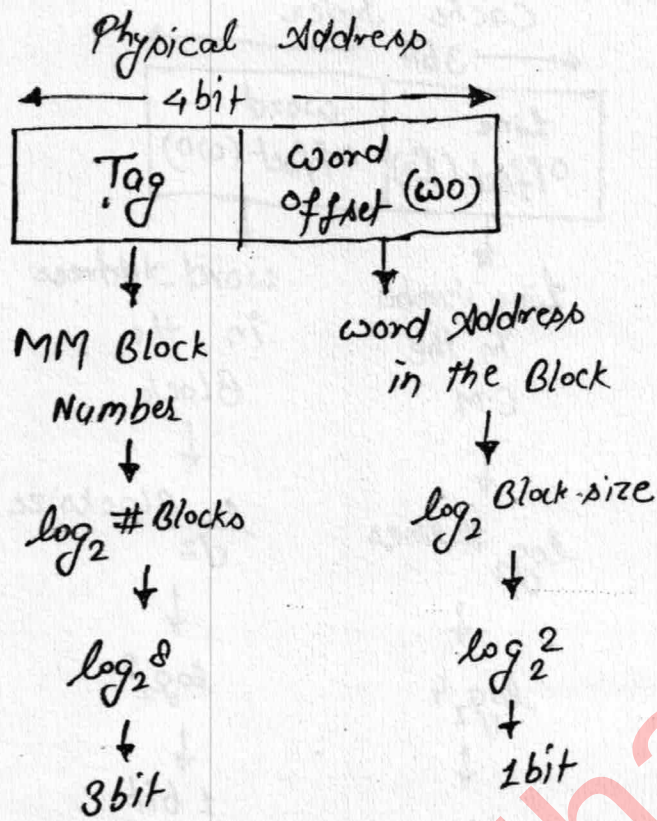


Consider 16B MM with
2B Block :-

Before orgn
16B → MM

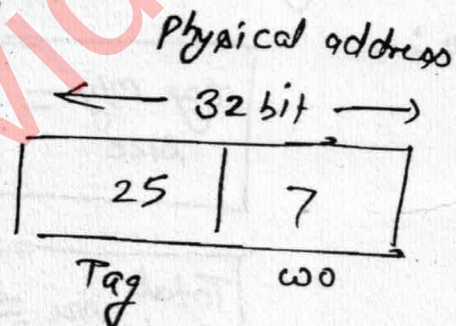
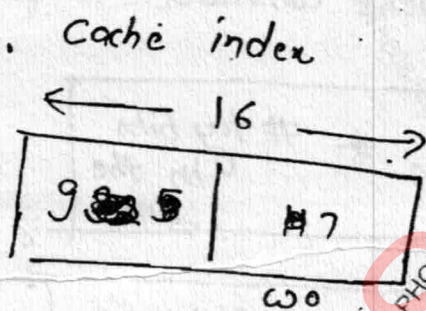
After orgn
 $\# \text{ Blocks} = \frac{16}{2} = 8$
 $8 * 2B = 16B$





Que:- Consider a system configuration with 64KB cache organized into a 32 word blocks. word size in the cpu is 32 bit. cpu generates 32 bit Physical address to access the data during the Program execution, How many lines and how many blocks are present in the respective memories. Show the address format :

$$\# \text{ Blocks in cache} = \underline{32}$$



$$\begin{aligned} \# \text{ lines} &= \frac{\text{Cache size}}{\text{Block size}} \\ &= \frac{64 \text{ KB}}{32 \text{ word}} \\ &= \frac{64 \text{ KB}}{32 \times 4 \text{ B}} \\ &= \frac{2^{16}}{2^7} \\ &= 2^9 \end{aligned}$$

$$\begin{aligned} \# \text{ Blocks} &= \frac{\text{MM size}}{\text{Block size}} \\ &= \frac{2^{32} \text{ B}}{32 \text{ w}} \\ &= \frac{2^{32} \text{ B}}{32 \times 4 \text{ B}} \\ &= \frac{2^{32}}{2^7} \\ &= 2^{25} \end{aligned}$$

mapping Techniques :-

:- The Process of transferring the data from main m/y to cache m/y is called as mapping.

* During the mapping Process Data Block is Copied into a cache memory along with a physical address.

→ To Hold the Physical address information, some extra Δ space is maintained in the cache line called as Tag-Space.

* Tag memory size in the cache controller is calculated as :

$$\text{tag m/y size} = \# \text{ lines in C.M.} * \# \text{ tag bits in the line}$$

$$\text{Total Cache size} = \text{Tag m/y size} + \text{Data m/y size}$$

$$\downarrow$$
$$\left\{ \# \text{ lines in C.M.} * \text{Block size} \right\}$$

* In the Cache Design, three kinds of mapping techniques are used to map the data name

as :

- 1) Direct mapping
- 2) Associative mapping
- 3) Set associative mapping