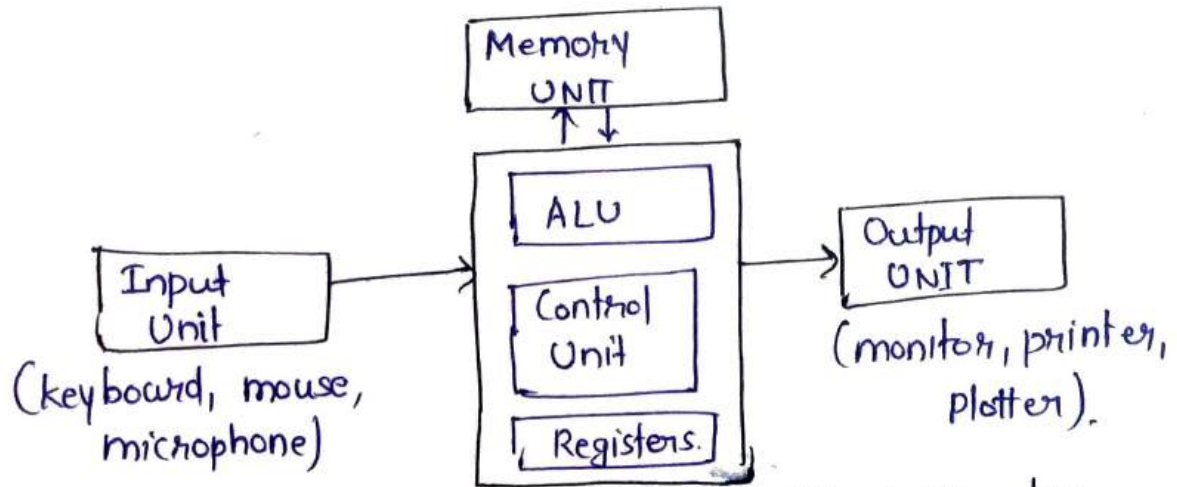


Computer Architecture & Organization

UNIT-I

Structure of Desktop computers



Memory — (RAM) CPU. ALU - Arithmetic & Logic Unit.

- Primary (fast, small, expensive) → store programs & active data.
- Secondary (large, slow, cheap)
↓
magnetic tapes, magnetic disks.

ALU —

- add, subtract, division, multiplication (arithmetic operation)
- (Logical operations) (AND, OR, Inverting).

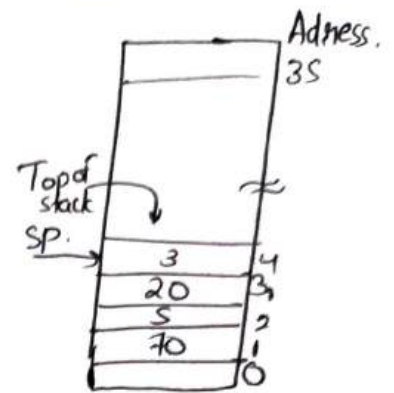
Control Unit —

- Control & Coordinate activities amongst all functional units.
- function —
 - fetch instruction from memory
 - decode (identify the operation)
 - execute instruction (sending control signals to corresponding unit).

Stack implementation

Register stack.
Memory stack.

It is organized as a collection of finite number of CPU registers.



- stack pointer holds the address of the register that is currently the top of stack.
- In above figure, 4 elements are stored in stack
- Data element 3 is top of stack, therefore content of $SP=4$.
- stack pointer is a 5 bit register, because there are 32 registers in stack ($2^5=32$)
- $SP=0$ means stack is empty
- When data element is pushed on the stack, SP is incremented.

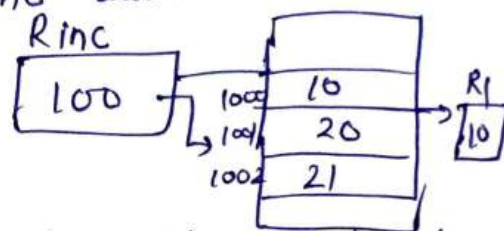
Memory stack

- It is implemented using computer memory.
- operation of memory stack is exactly similar to register stack.
- Number of registers in the CPU is limited & it restricts the size of stack.
- Using memory, size of stack can be extended upto size of memory.

— Or Address of the operand is in a register whose value is incremental/decremented after fetching the operand from the address.

Ex: Load (R inc)+, R₁

$$① R_1 \leftarrow M[R_{inc}]$$



Load the contents of the memory address stored in register R inc.

② $[R_{inc}]++$ Or $R_{inc} = R_{inc} + 1$
 ↳ increment the ~~control~~ contents of R inc.
 (memory address)

OR INR R₁ will increment the Register R₁

— When the address stored in the register refers to a table of data in memory, it is necessary to increment/decrement the register after every access to the table.

— This can be achieved using increment or decrement instructions.
 DCR R₂ - will decrement the register R₂

Auto decrement.

Ex Load (R dec)-, R₁

$$② R_1 \leftarrow M[R_{dec}]$$

$[R_{dec}]_-$ Or $[R_{dec}] \leftarrow [R_{dec}] - 1$ [Before access of execution of instruction]

★ This mode is specify useful when we want to access a table of data.

① Arithmetic instructions -

Name	Mnemonic	Description
Increment	INC	Add 1 to the value stored in a register or memory word
Decrement	DEC	Sub 1 from the value stored in a register or memory word.
Add	ADD	ADD 2 numbers (binary, flt pt, deci)
Subtract	SUB	Subtract 2 no.
Multiply	MUL	Multiply 2 no
Divide	DIV	DIV 2 no.
Add with carry	ADDC	Performs the addition of 2 operand plus the value of the carry from previous computation.
Subtract with borrow	SUBB	Subtracts 2 words & a borrow which may have resulted from previous subtract operation
Negative 2's Complement	NEG	From the 2's complement of a number effectively.

② Logical & Bit Manipulation instructions. -

Name	Mnemonic	Description
Clear	CLR	0's transferred to the destination to clear selected bits.
Complement	COM	
AND	AND	
OR	OR	
Exclusive OR	XOR	
Clear carry	CLRC	
Set carry	SETC	
Complement carry	COMC	
Enable Interrupt	EI	- Enable interrupt flip flop that control is the ! interrupt facility is enabled.
Disable interrupt	DI	- Flip flop that controls the interrupt facility is disabled.

- The purpose of micro program sequencer is to present an address to the control memory so that a microinstruction may be read & executed.
- Here multiplexer selects an address from 4 sources & routes it into control address register.
- The output from CAR provides the address for the control memory.
- The contents of CAR are incremented & applied to the multiplexer & to the stack register file.
- - The register selected in the stack is determined by the stack pointer.
- Inputs I_1, I_2, I_3 & T (test) derived from CD & BR fields of microinstruction specify the operation for the sequencer.
- They specify the input source to the multiplexer also generate push & pop signals required for stack operation.

Typical Sequencer operations are

- Increment
- Branch
- Jump
- Call & return from Subroutine
- load an external address
- push or pop the stack
- & other address sequencing operations.
- The stack pointer is a 3 bit register & it gives the address of stack register file consists of ($2^3 = 8$) eight registers.

- ③ If it is, then the instruction is fetched from the cache - a very fast process (cache-hit)
- ④ If not, the CPU has to fetch next instruction from main memory - a much slower process (cache miss)

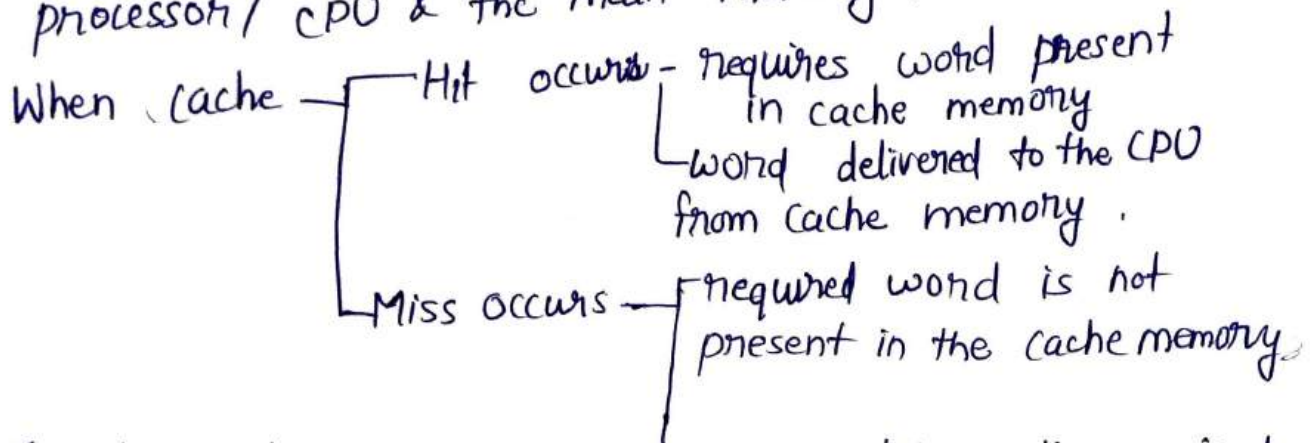
Hit Ratio - The performance of cache memory is measured in terms of quantity called hit ratio.

- When the CPU refers to memory & finds the word in cache it is said to provide a hit.
- If the word is not found in cache, it is in main memory then it counts as a miss.
- The ratio of the number of hits to the total CPU references to memory is called hit ratio.

$$\text{Hit ratio} = \frac{\text{Hit}}{\text{Hit} + \text{miss}}$$

CACHE MAPPING TECHNIQUES.

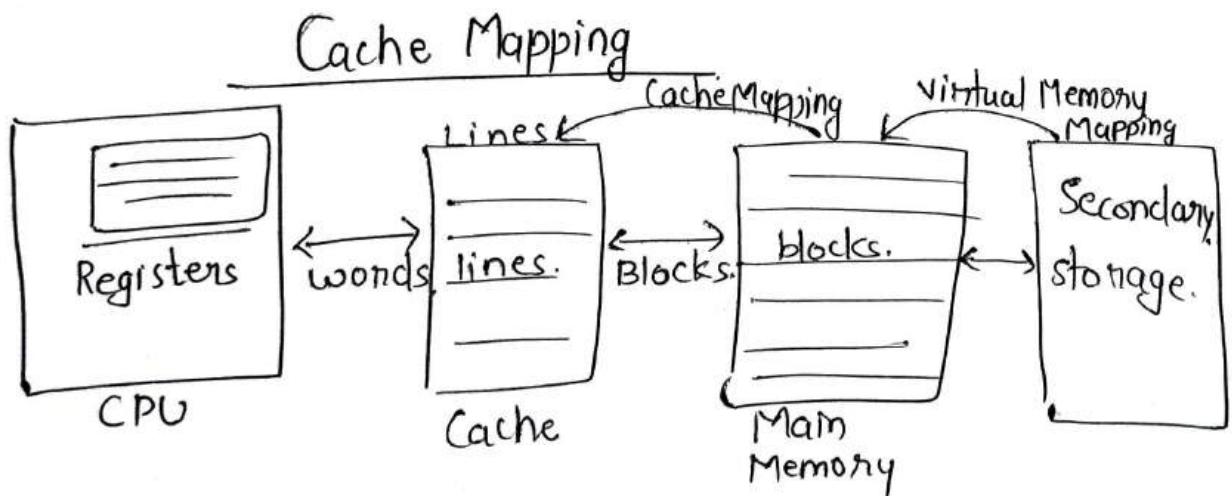
Cache memory bridges the speed mismatch b/w processor / CPU & the main memory.



So when cache miss occurs, the page containing the required word has to be transferred from the main memory to Cache memory.

- The transformation of data from main memory to Cache memory is called as mapping process.
This mapping is performed using cache mapping techniques.

Defⁿ Cache mapping is a technique by which the contents of main memory are brought into the cache memory.



Mapping Process.

During cache mapping, block of main memory is simply copied to the cache & the block is not actually brought from the main memory.

There are three types of cache Mapping Techniques.

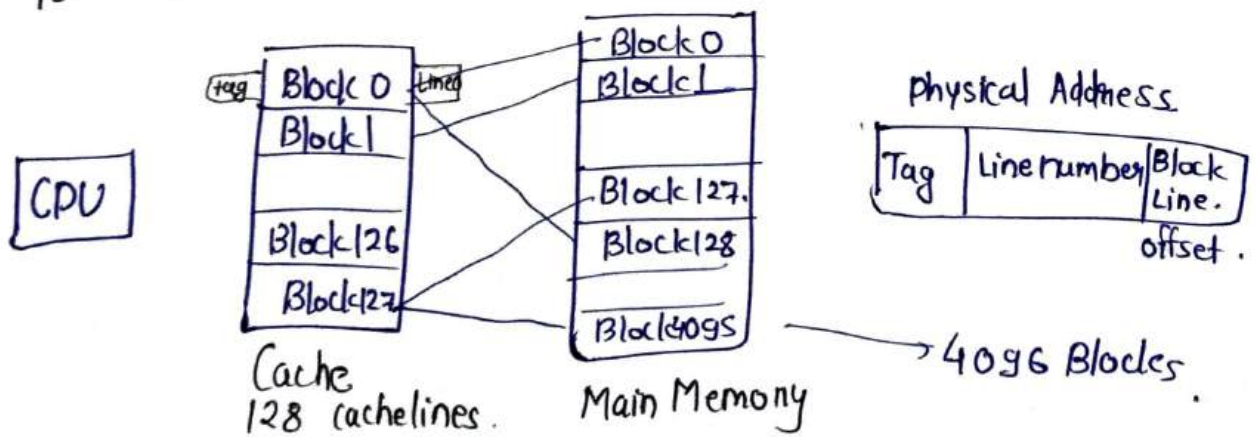
- ① Direct Mapping.
- ② Fully Associative Mapping.
- ③ k-way set associative Mapping.

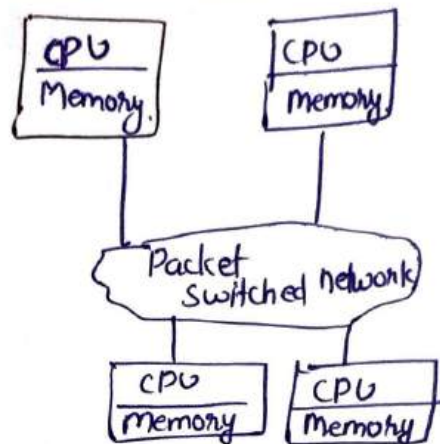
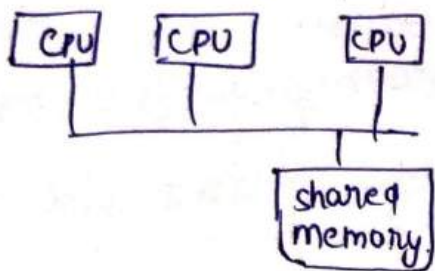
① Direct Mapping - In direct mapping, a particular block of main memory can map only to a particular line of the cache.

- The line number of cache to which a particular memory block can map is given by -

$$\text{Cache line number} = \frac{(\text{main memory Block Address}) \bmod (\text{Number of lines in Cache})}{\text{Number of lines in Cache}}$$

Ex. Consider cache memory is divided into n number of lines then, block 'j' of main memory can map to the line number $(j \bmod n)$ only of the cache.





Difference between Tightly Coupled and Loosely Coupled Multiprocessor.

Tightly Coupled

- ① Tightly coupled multiprocessor has shared memory.
- ② The degree of coupling between processors is high.
- ③ In tightly coupled multiprocessors CPU's connected in close communication.
- ④ In this system, CPU's share computer Bus, memory & I/O devices.
- ⑤ It is efficient when high degree of interaction between processes & for high speed & real time processing.
- ⑥ Memory conflict occurs due to use of shared memory.

Loosely Coupled

- ① Loosely coupled multiprocessor has distributed memory.
- ② The degree of coupling b/w processors is low.
- ③ In loosely coupled multiprocessors CPU/systems located at different locations.
- ④ In this system there is no such sharing.
- ⑤ It is efficient when tasks running on different processors has minimal interaction b/w them (parallel application).
- ⑥ No memory conflict occurs.

⑦ Processors communicate with each other using shared memory.

⑧ Data rate is high

⑨ More expensive

⑩ Compact in size

⑪ Less scalable
- limited number of CPU can be added

⑫ Less fault tolerant
- If shared memory corrupted then causing all processor to fail.

⑬ Less complex

⑭ Portable

⑧ Processors communicate by using message passing Techniques.

⑨ Data rate is low

⑩ Less expensive.

⑪ Larger in size.

⑫ Scalability is high.
- more no. of CPU's can be added to improve the system performance.

⑬ more fault tolerant
- A fault in a single system/module does not lead to a complete system breakdown.

⑭ Complex due to additional hardware is required to provide communication between individual processors.

⑮ Less portable.

Now P₁ makes an exit from the CS, it will execute "signal" operation $S=1$.

Then one of the processes, looping in "while" statement of "wait" operation $\rightarrow S=0$ enter its CS.

As "wait" operation is to be executed automatically, so other waiting processes will find the $S=0$ & ~~can~~ continue to loop in the "while" loop.

So, at a time only one of the cooperating process can enter CS, subject to satisfaction of the condition that "wait" operation is executed automatically.

★ So the requirement of mutual exclusion is met.

(ii) Progress - $S=1$ - no process is executed in CS.

One of the waiting processes looping in the "while" loop of "wait" operation will find the $S=1$, exit from the "while" loop decrement the semaphore to 0 & enter CS.

So if no process is executing in CS & some are waiting to enter, then one of the waiting processes will enter its CS immediately. Thus the requirement of progress is met.

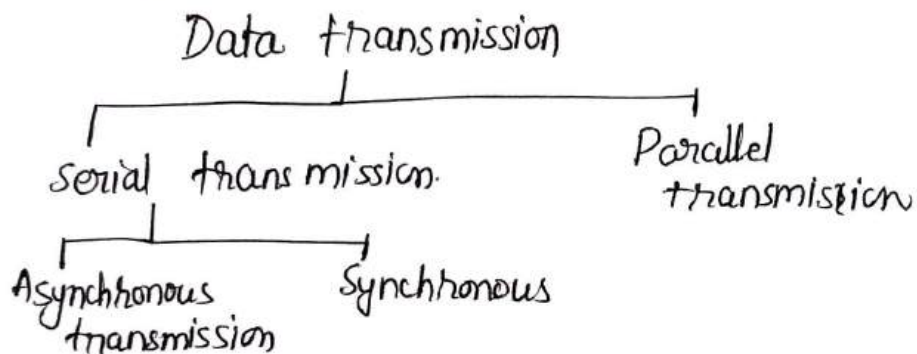
(iii) Bounded Waiting -

Arbitrary, one of the waiting processes, will get entry into its critical section, when a cooperating process, executing in its critical section exists.

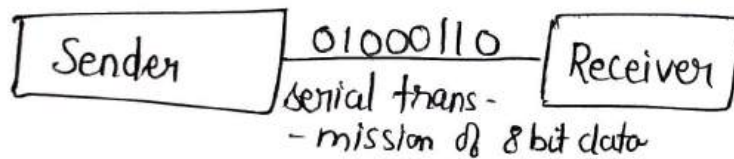
DATA TRANSMISSION

Data transmission refers to the process of transferring data b/w 2 or more digital devices.

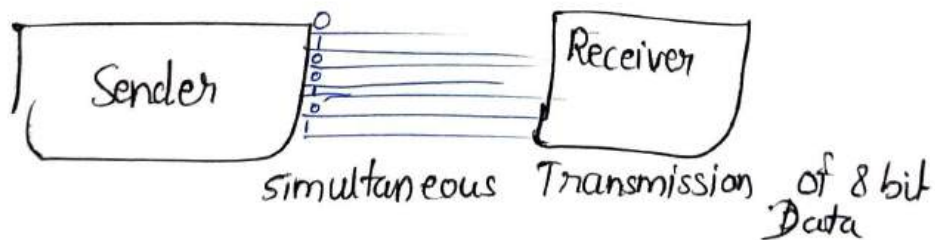
- Data transmitted from 1 device to another in analog or digital format.
- Data is transferred in the form of bits b/w 2 or more digital devices.



Serial transmission



Parallel transmission.



- | Serial data transmission | Parallel Data transfer |
|---|--|
| ① Serial data transmission sends data bits one after another over a single channel. | ② Parallel data transmission sends multiple data bits at the same time over multiple channels. |