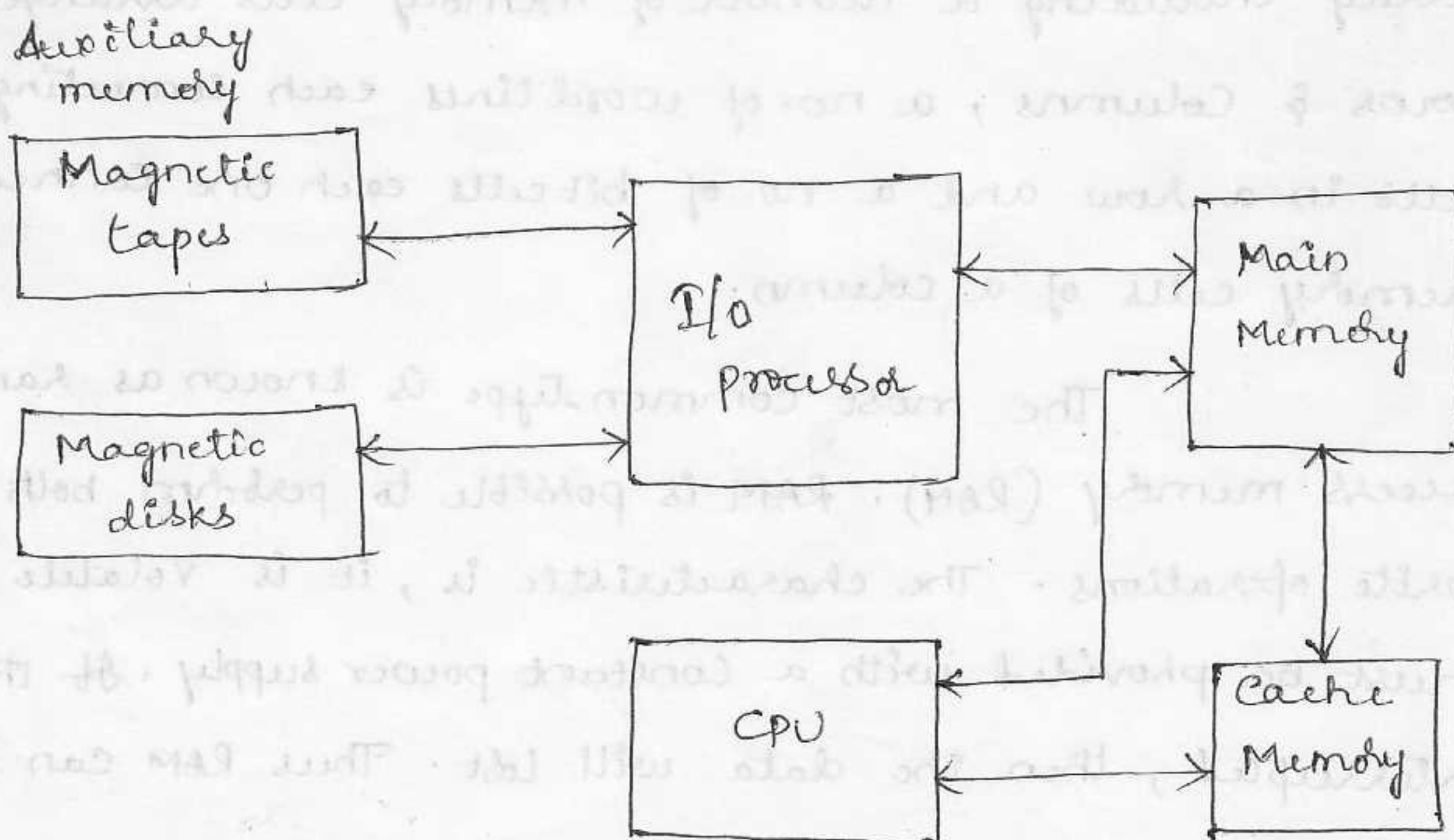


MEMORY ORGANIZATIONMemory Hierarchy :-

The memory unit that communicates directly with the CPU is called the main memory. Devices that provide backup storage are called auxiliary memory. The most common auxiliary devices used in computer systems are magnetic disks and tapes. They are used for storing system programs, large databases and other back up information. Only programs and data currently needed by the processor reside in main memory and all other information is stored in auxiliary memory and transferred to main memory when it is needed.

The below fig. shows the components in a memory hierarchy.



The Memory hierarchy system consists of all storage devices employed in a computer system from the slow but high-capacity auxiliary memory to a relatively faster main memory and

187
Cache memory accessible to the high-speed processing logic.

At the bottom of the hierarchy are the relatively slow magnetic tapes used to store removable files. Next are the magnetic disks are used as backup storage. The main memory occupies a central position be able to communicate directly with the CPU and with auxiliary memory devices through an I/O processor.

When programs not residing in main memory are needed by the CPU, they are brought in from auxiliary memory.

A special very-high-speed memory called a cache. It is used to increase the speed of processing by making current programs and data available to the CPU at a rapid rate.

Semiconductor RAM Memory :-

The semiconductor memory includes a memory cell array including a number of memory cells arranged in a rows & columns, a no. of word lines each connecting all memory cells in a row and a no. of bit cells each one connecting all memory cells of a column.

The most common type is known as random access memory (RAM). RAM is possible to perform both read and write operations. The characteristic is, it is volatile. A RAM must be provided with a constant power supply. If the power is interrupted, then the data will lost. Thus RAM can be used only as temporary storage.

RAM technology has divided into two technologies

- (i) Static RAM
- (ii) Dynamic RAM.

Static RAM :-

In a static RAM, binary values are stored as using traditional flip-flop logic gate configurations. A SRAM will hold its data as long as power is supplied to it. SRAM is used as cache memory.

Dynamic RAM :-

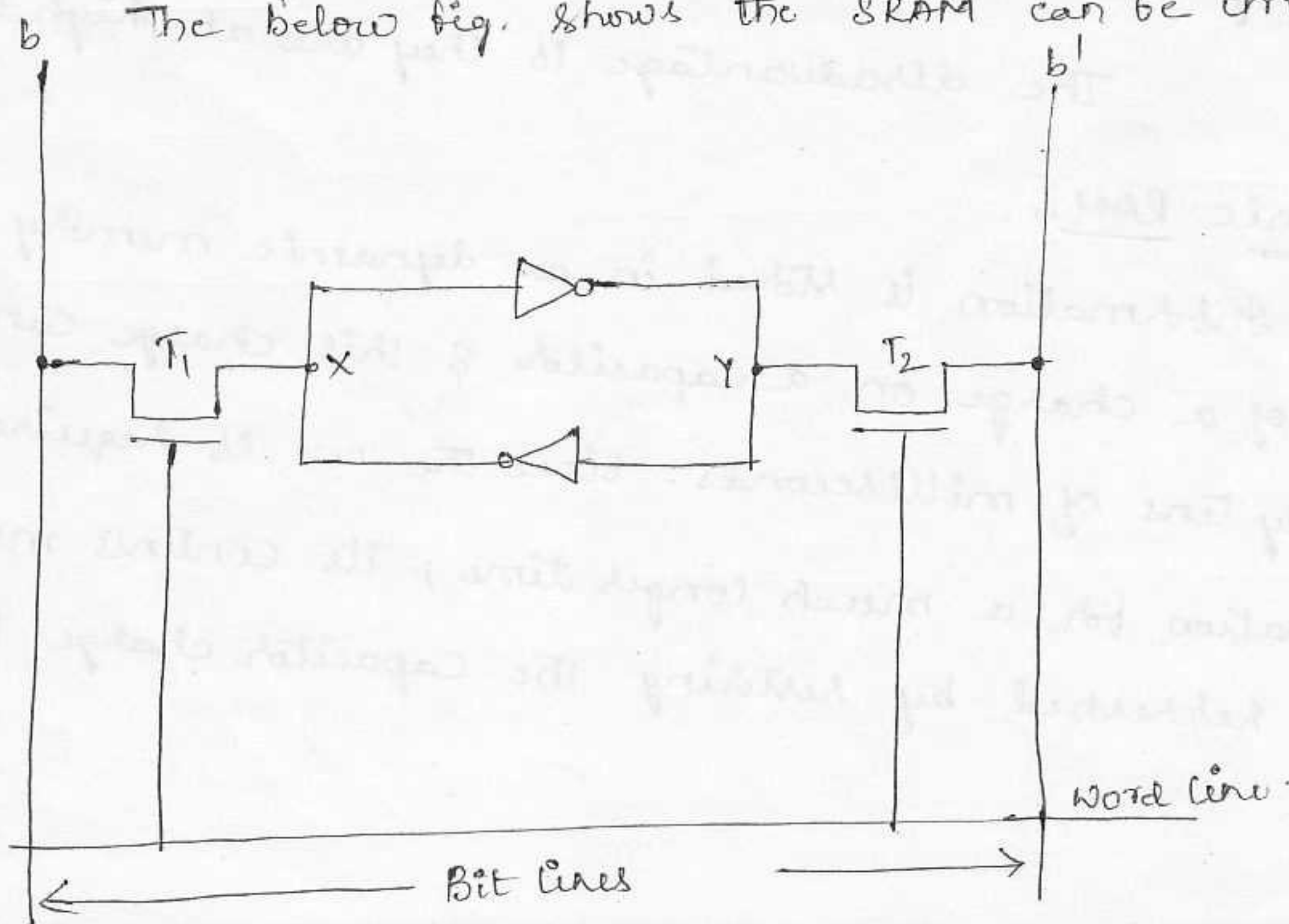
A DRAM is made with cells that store data as charge on capacitors.

DRAM's are among the densest VLSI circuit in terms of transistors per chip.

Both SRAM's and DRAM's are volatile, that is the stored information is lost when the power is removed. DRAM is used as main memory. A dynamic memory cell is simpler and hence smaller than a static memory cell. Thus, a dynamic RAM is more dense and less expensive than SRAM. But SRAM's are generally faster than DRAM's.

There is another type of memory is known as Read-only memory (ROM).

The below fig. shows the SRAM can be implemented.



In the bit; the two inverters are cross-connected to form a latch. The latch is connected to two bitlines by transistors T_1 & T_2 . These transistors act as switches that can be opened or closed under the control of the word line. When the word line is at ground level, the transistors are turned off and the latch retains its state. For eg, let us assume that the cell is state 1 if the logic value at point X is 1 and at point Y is 0.

Read operation:

To read in SRAM cell, the wordline is activated to close switches T_1 & T_2 . If the cell is in state 1, the signal on bitline 'b' is high & the signal on bitline 'b'' is low.

Write operation:-

The state of the cell is set by placing the appropriate value on bitline 'b' and its complement on b' and then activating the wordline.

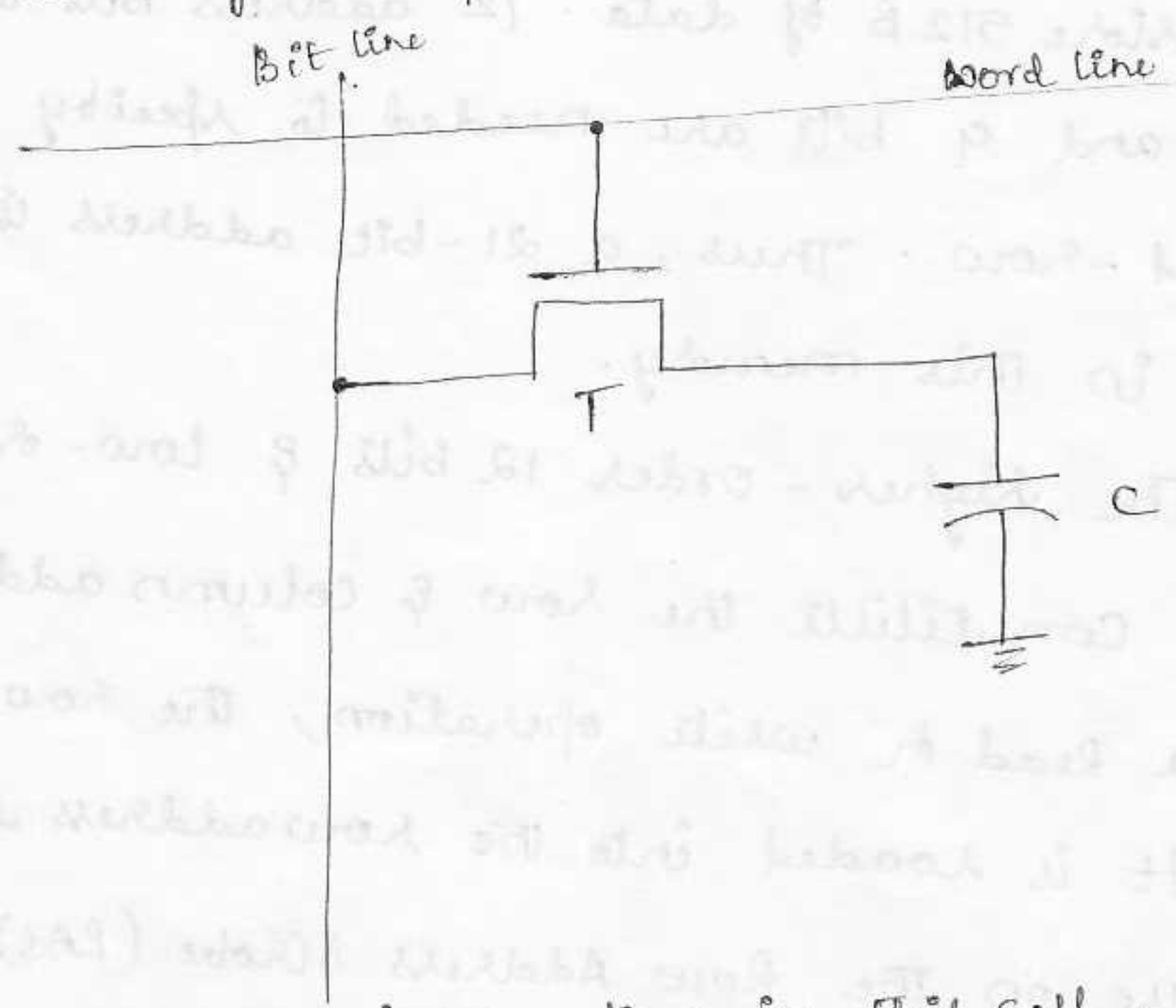
A major advantage of SRAM's is their very low consumption because current flows in the cell only when the cell is being accessed and also they are fast.

The disadvantage is they are at high cost.

Dynamic RAM's:

Information is stored in a dynamic memory cell in the form of a charge on a capacitor & this charge can be maintained for only tens of milliseconds. Since the cell is required to store information for a much longer time, its contents must be periodically refreshed by restoring the capacitor charge to its full value.

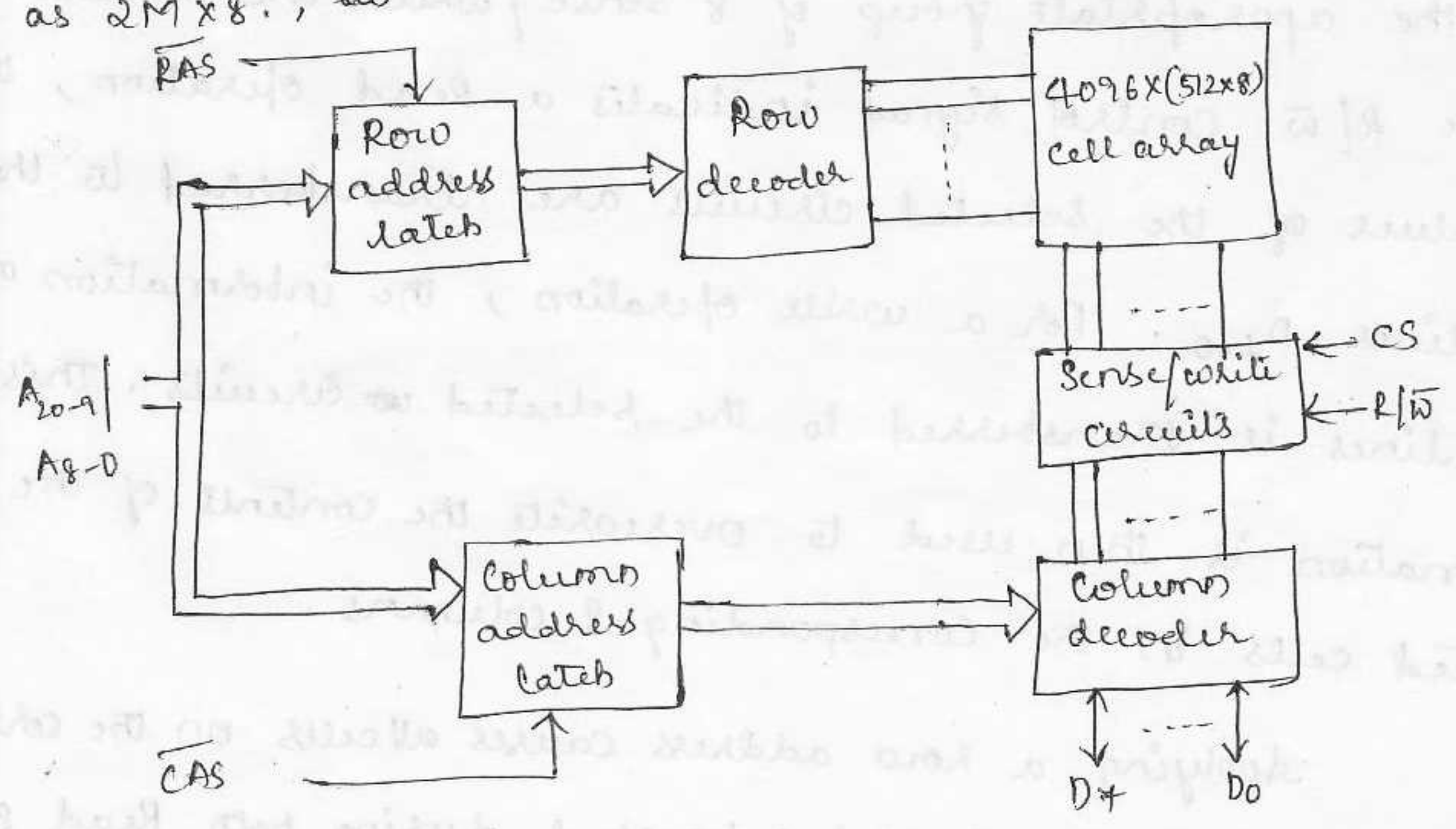
The below fig. shows the dynamic memory cell that consists of a capacitor C and a transistor 'T'.



In order to store information in this cell, the transistor T is turned on and an appropriate voltage is applied to the bit line. This causes a known amount of charge to be stored in the capacitor. After the transistor is turned off, the capacitor begins to discharge.

Asynchronous DRAM's :-

The below fig. shows the 16-Megabit DRAM chip, configured as $2M \times 8$.



The cells are organized in the form of a $4K \times 4K$ array. The 4096 cells in each row are divided into 512 groups of 8, so that each row can store 512 B of data. 12 address bits are needed to select a row and 9 bits are needed to specify a group of 8 bits in the selected row. Thus, a 21-bit address is needed to access a byte in this memory.

The higher-order 12 bits & low-order of 9 bits of the address constitute the row & column address of a byte.

During a read or write operation, the row address is applied first. It is loaded into the row address latch in response to a signal pulse on the Row Address Strobe (RAS) in the chip. Then a Read operation is initiated, in which all cells on the selected row are read & refreshed.

After the row address is loaded, the column is applied to the address pins and loaded into the column address latch under the control of the column Address Strobe (CAS) signal.

The information in the latches are decoded and the appropriate group of 8 sense/write are selected. If the R/W control signal indicates a Read operation, the O/p values of the selected circuits are transferred to the data lines D_{7-0} . For a write operation, the information on the D_{7-0} lines is transferred to the selected circuits. This information is then used to overwrite the contents of the selected cells in the corresponding 8 columns.

Applying a row address causes all cells on the corresponding row to be read and refreshed during both Read &

19.2

while operations. To ensure that the contents of a DRAM are maintained, each row of cells must be accessed periodically. A refresh counter performs this function automatically.

The timing of the memory device is controlled asynchronously. A specialized memory controller circuit provides the necessary control signals, RAS & CAS signals. The processor must take into account the delay in the response of the memory. Such memories are referred to as asynchronous DRAM's.

Advantages :-

- 1) It has high density & low cost, these are widely used.
- 2) Available chips are of size 1M to 256 Mbits & larger chips are developed.
- 3) A DRAM chip is organized to read or write a number of bits in parallel.
- 4) It provides flexibility in designing memory systems.

Fast page Mode :-

The contents of all 4096 cells in the selected row are sensed, but only 8 bits are placed in the data bus lines D_7-0 . This byte is selected by the column address bits A_8-0 .

A simple modification can make it possible is, to access the other bytes in the same row without having to reselect the row. A latch can be added at the O/p of the sense

amplifier in each column.

The application of a row address will load the latches corresponding to all bits in the selected row. Then, it is only necessary to apply different column address to place the different bytes on the data lines.

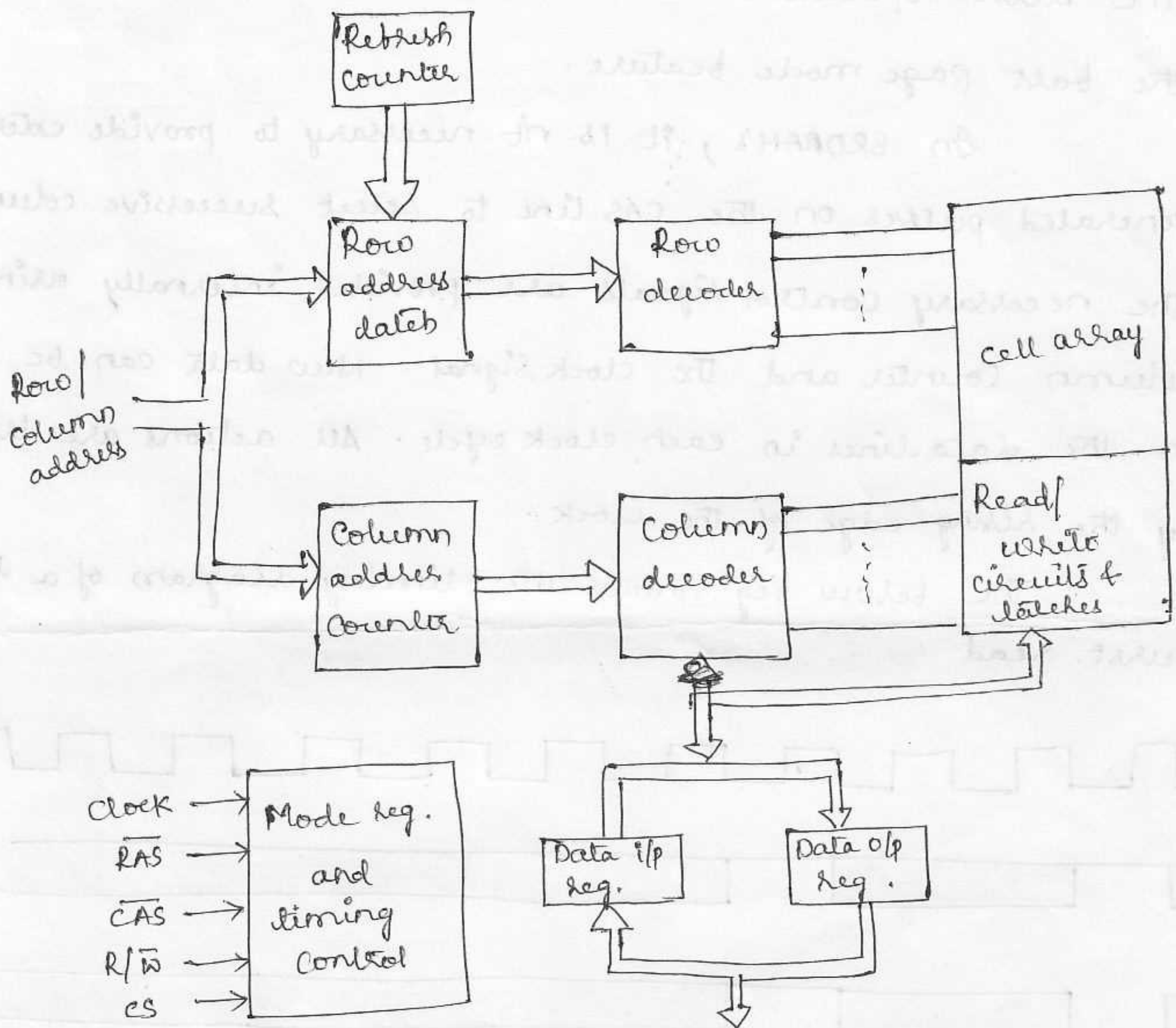
The most useful arrangement is to transfer bytes in sequential order, which is applied by applying a consecutive sequence of column address under the control of successive CAS signals.

This scheme allows transferring a block of data at a much faster rate than can be achieved for transfers involving random addresses. The block transfer capability is referred to as 'burst page mode' features.

Synchronous DRAM's :-

The operation is directly synchronized with a clock signal, such memories are known as Synchronous DRAM's.

The below figure shows the structure SDRAM.



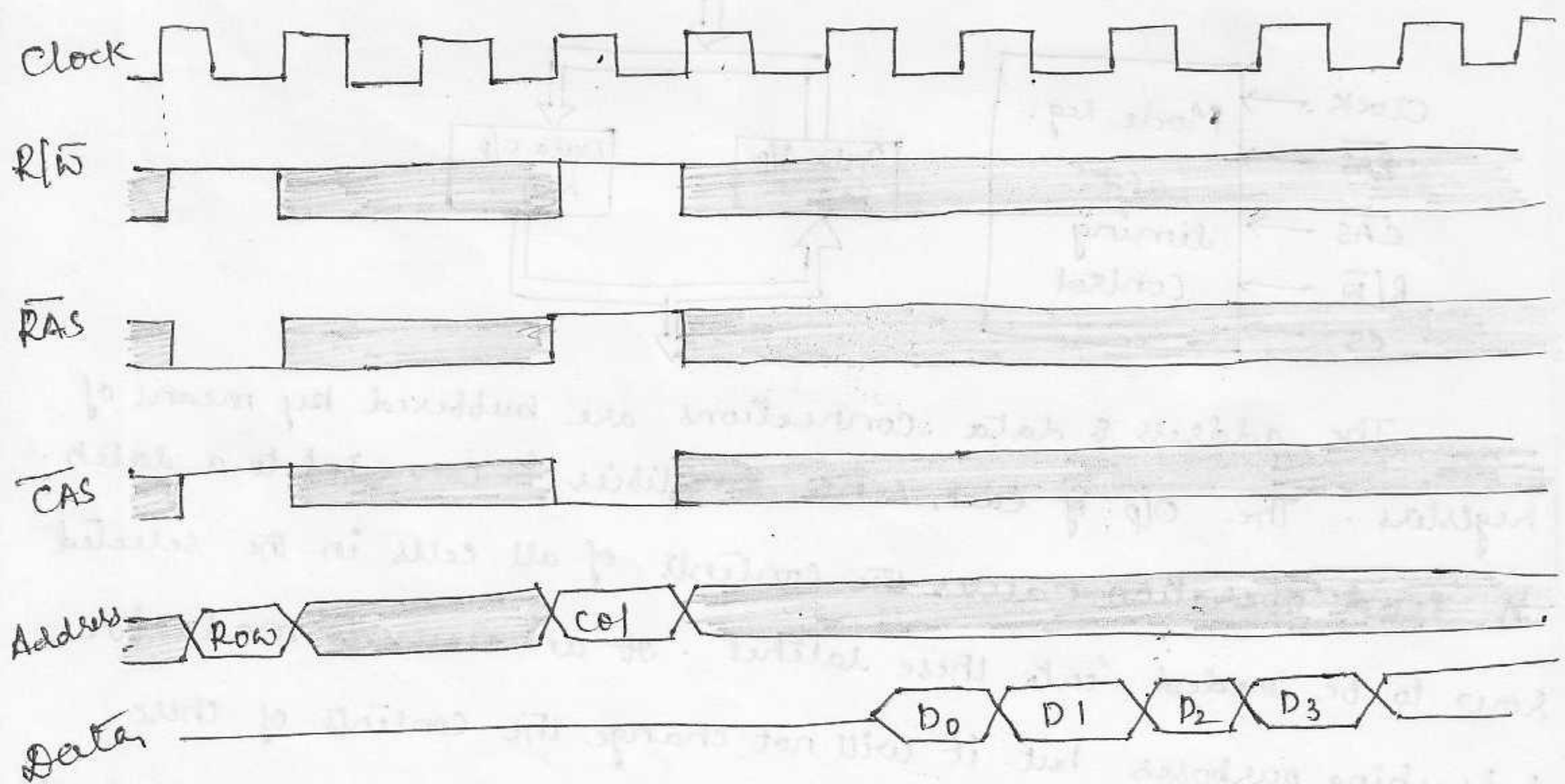
The address & data connections are buffered by means of registers. The o/p of each sense amplifier is connected to a latch. A read operation causes the contents of all cells in the selected row to be loaded into these latches. At an access is made for refreshing purposes but it will not change the contents of these latches and it will refresh the contents of the cells. Data held in the latches that correspond to the selected columns are transferred into the data o/p reg.

SRAM's have several different modes of operation, which can be selected by writing control information into 'mode' register.

For eg, burst operations of different lengths can be specified. The burst operations use the block transfer capability as in the burst page mode feature.

In SRDRAM's, it is not necessary to provide externally generated pulses on the CAS line to select successive columns. The necessary control signals are provided internally using a column counter and the clock signal. New data can be placed on the data lines in each clock cycle. All actions are triggered by the rising edge of the clock.

The below fig shows the timing diagram of a typical burst read.



First, the row address is latched under control of the RAS signal. The memory typically takes 2 or 3 clock cycles to activate the selected row. Then, the column address is

196

Latched under control of $\overline{\text{CAS}}$ signal. After a delay of one clock cycle, the first set of data bits is placed on the data lines.

The SDRAM automatically increments the column address to access the next three sets of bits in the selected row, which are placed on the data lines in the next 3 clock cycles.

SDRAMs have built-in refresh circuitry. A part of this circuitry is a refresh counter, which provides the addresses of the rows that are selected for refreshing.

The Commercial SDRAM's can be used with clock speeds above 100 MHz.

~~Organization of DRAM's :-~~

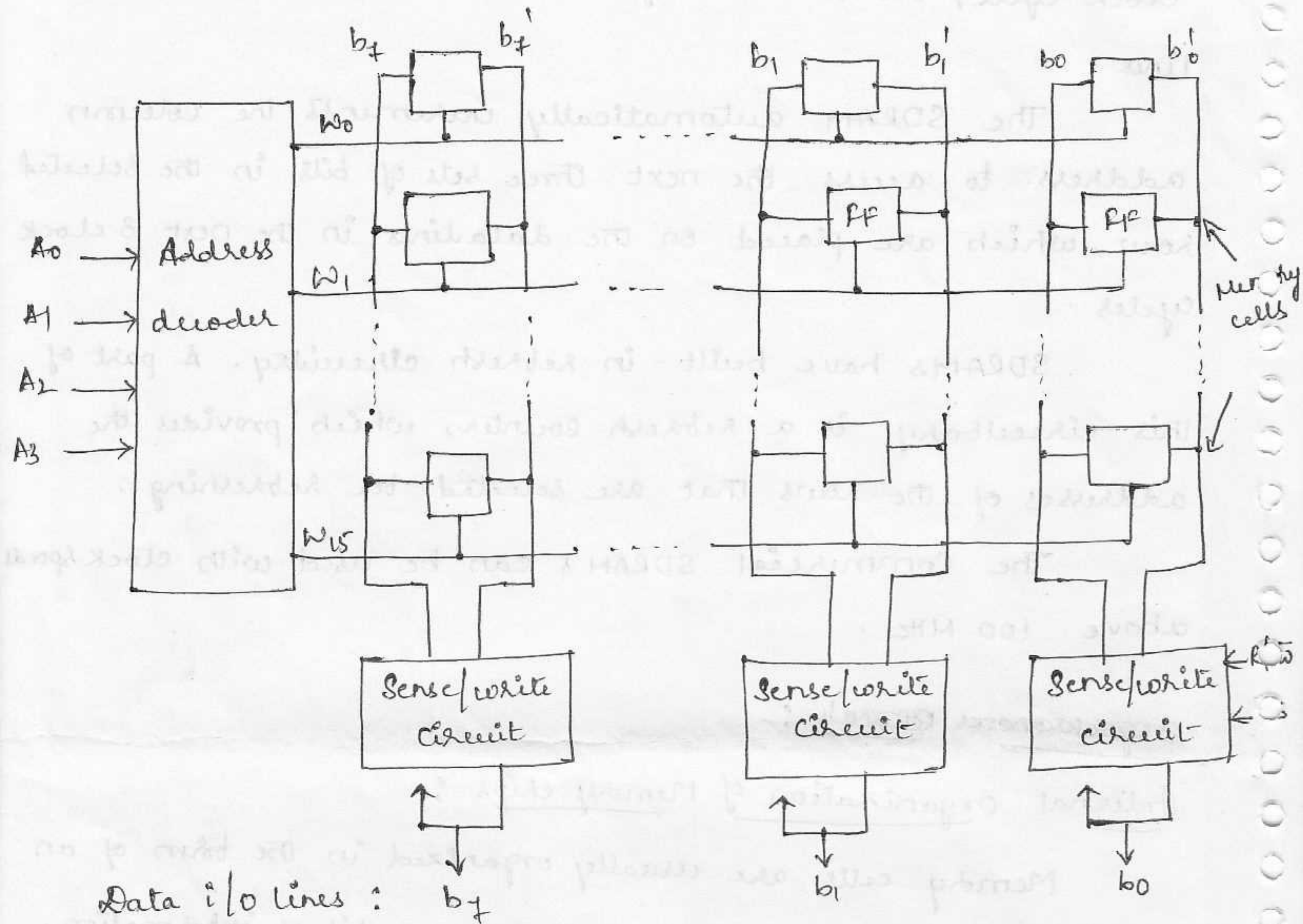
Internal Organization of Memory chips :-

Memory cells are usually organized in the form of an array, each cell is capable of storing one bit of information.

The fig. is an example of a very small memory chip consisting of 16 words of 8 bits each.

Each row of cells constitutes a memory word and all cells of a row are connected to a common line is referred to as the word line, which is driven by the address decoder on the chip. The cells in each column are connected to a sense/write circuit by two bit lines.

The sense/write circuits are connected to the data i/o lines of the chip. During read operation, these circuits sense, or read the information stored in the cells selected by a



Data i/o lines : b_7, b_1, b_0

Word line and transmit this information to the o/p data lines. During a write operation, the sense/write circuit receive the information and store the cells of the selected ~~the~~ word.

The data i/p & the data o/p of each sense/write circuit are connected to a single bidirectional data line that can be connected to the data bus of a computer.

Two control lines, $R/\bar{W}$ & $\bar{CS}$ are provided in addition to the address & data lines. The $R/\bar{W}$ input specifies the required operation & the $\bar{CS}$ (chip select) i/p selects a given chip in a multi chip memory system.

PROM :-

For small quantities, we are going to use a type of a ROM called programmable Read only Memory.

In this, the programmability is achieved by inserting a fuse at point p (shown in fig.). Before it is programmed, the memory contains all 0's.

The user can insert 1's at the required locations by burning out the fuses at these locations using high-current pulses.

PROM's provide flexibility & Convenience

PROM's provide faster and considerably less expensive because they can be programmed directly by the user. The ROM & PROM are irreversible.

EPROM :-

This is an type of ROM, which allows the stored data to be erased and new data to be loaded. Such an erasable, reprogrammable ROM is usually called an EPROM.

EPROM's are capable of retaining stored information for a long time, they can be used in place of ROM's.

The EPROM & ROM has similar structure. In EPROM cell, the connection to ground is always made at point 'p' and a special transistor is used, which has the ability to function as either as a normal transistor or as a disabled transistor that is always turned off. i.e., the transistor can be programmed to behave as a permanently open switch.

The advantage is that their contents can be erased and reprogrammed. The disadvantage is that a chip is removed physically from the circuit for reprogramming and the entire contents is erased by exposure to U.V. light.

EEPROM :-

The another version of erasable PROM's that can be both programmed and erased electrically. Such chips are called EEPROM's. It is possible to erase the cell contents selectively.

The disadv. is that different voltages are needed for erasing, writing and reading the stored data.

Flash Memory :-

Flash Memory is intermediate b/w EPROM & EEPROM. Like EEPROM, flash memory uses an electrical erasing technology. An entire flash memory can be erased in one or a few seconds which is much faster than EPROM.

Flash Memory uses only one transistor per bit and so achieves the high density.

Cache memories :-

The speed of the main memory is very low, because for good performance, the processor cannot ^{spend} much of its time waiting to access instructions and data in main memory.

An efficient solution is to use a fast cache memory which essentially makes the main memory appear to the processor to be faster.

The effectiveness of cache mechanism is based on a property of computer programs called 'locality of reference'. Analysis of programs show that most of their execution time is spent on routines in which many instructions are executed repeatedly.

The instructions in localized areas of the program are

executed repeatedly during some time period and the remainder of the program is accessed relatively infrequently. This is referred to as 'locality of reference'. This is divided into two types 1) Temporal. 2) Spatial.

The temporal means that a recently executed instruction is likely to be executed again.

The spatial ~~as~~ means that instructions in close proximity to a recently executed instruction (Based on 'instruction address')

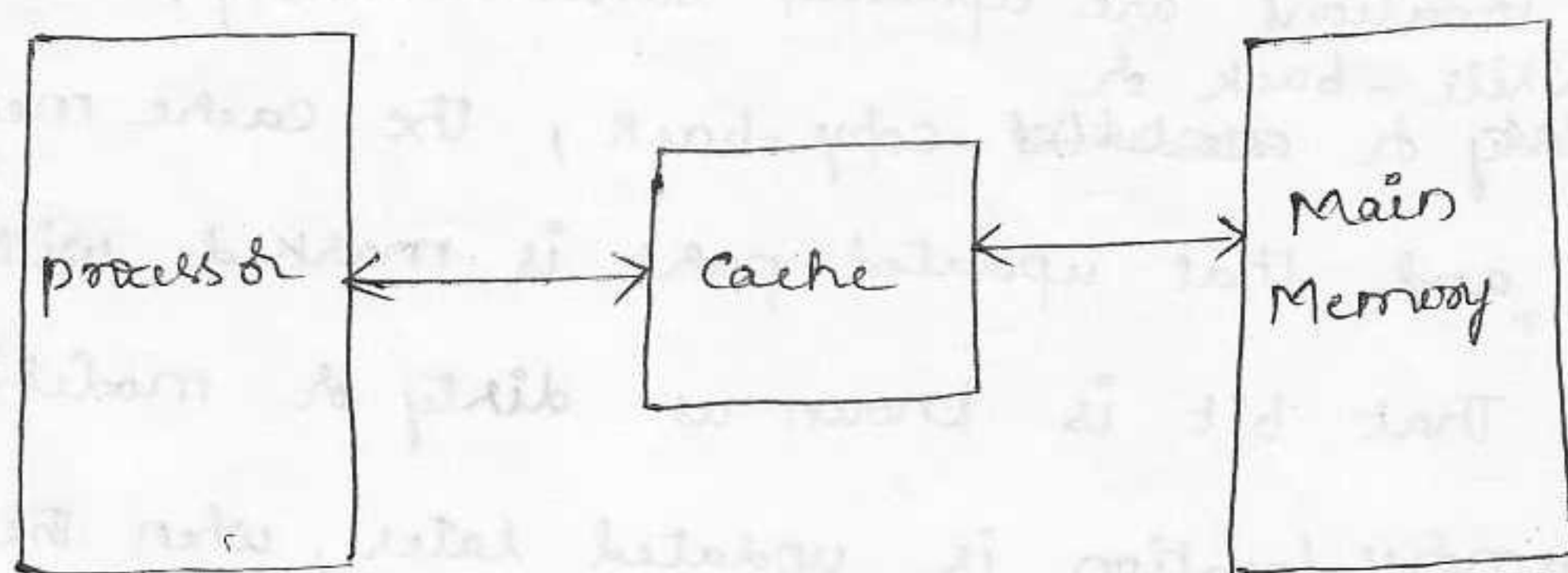
If the active segments of a program can be placed in a fast cache memory, then the total execution time is reduced.

The memory control circuitary is designed to take adv. of the property of locality of reference. The temporal aspect of the locality of reference, that whenever an information item is first needed, this item should be brought into cache where it is needed.

The Spatial aspect suggests that instead of fetching one item from the main memory to the cache, it is useful to fetch several items that reside at adjacent addresses.

The term 'block' refers to a set of contiguous address locations of some size. Another term is refer to a cache block is 'cache line'.

The below fig. shows the use of a cache memory.



When a Read Request is received from the processor, the contents of a block of memory words containing the location specified are transferred into the cache one word at a time. When the program references any of the locations in this block, the desired contents are read directly from the cache. The cache memory can store a reasonable number of blocks at any given time, but this number is small compared to the total number of blocks in the main memory.

The correspondence b/w the main memory blocks and those in the cache is specified by a mapping function.

When the cache is full and a memory word that is not in the cache is referenced, the cache control h/w must decide which block should be removed to create space for the new block that contains the referenced word. The collection of rules for making this decision is "replacement algorithm".

The cache control circuitry determines whether the requested word currently exists in the cache. The Read or Write operation is performed on the appropriate cache location. In Read operation, the main memory is not involved.

The write operation may be proceed into two ways

- 1) Write-through protocol
- 2) ~~dirty or modified bit~~ ^{Write-back or copy-back}

In write-through protocol, both cache memory and the main memory locations are updated simultaneously.

In ~~dirty or modified~~ ^{Write-back or} copy-back, the cache memory is only updated and that updated part is marked with its associated flag bit. That bit is known as dirty or modified bit. The main memory location is updated later, when the block containing

204

this marked word is to be removed from the cache to make room for a new block.

When the addressed word in a Read operation is not in the cache, a "read miss" occurs. The block of words that contains the requested word copies from the main memory into the cache. After the entire block is loaded into the cache, the particular word requested is forwarded to the processor.

The latter approach, which is called load-through or early restart, to this, it reduces the processor's waiting period, but it is more complex circuitary.

During a write operation, if the addressed word is not in the cache, a "write miss" occurs. Then, the write-through is used, the information is written directly into the main memory.

In the case of write-back protocol, the block containing the addressed word is first brought into the cache, and then the desired word in the cache is overwritten with the new information.

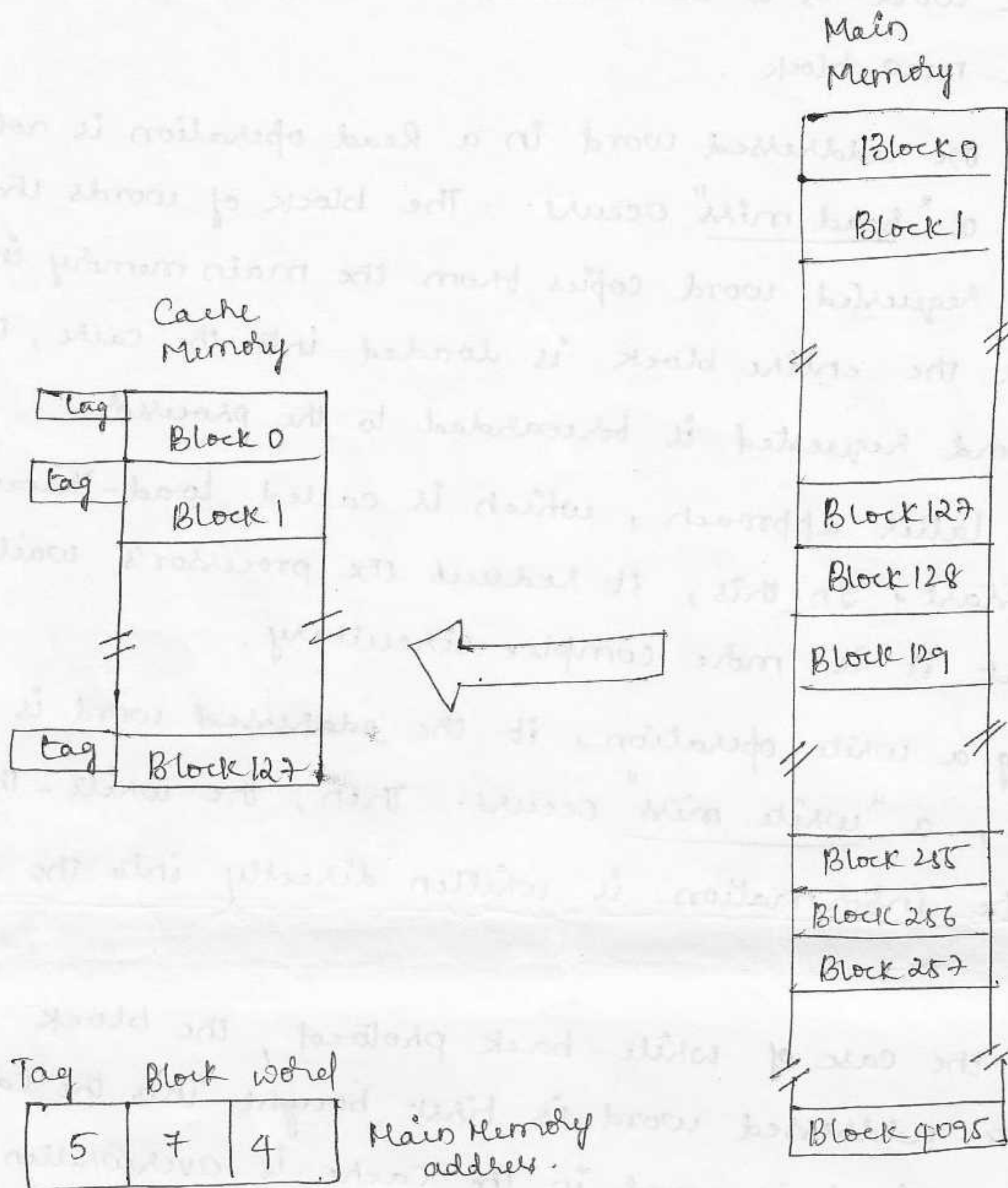
Mapping Functions :-

Direct Mapping :-

The simplest way to determine cache locations in which to store memory blocks is the direct-mapping technique.

In this technique, block j of the main memory maps onto block j modulo 128 of the cache, as depicted & shown in fig.

One of the main memory blocks 0, 128, 256 ... are loaded into ~~main memory~~ cache, it is stored in cache block 0. 1, 129, 257 ... are stored in cache block 1 and so on.



Since more than one memory block is mapped onto a given Cache block position, Contention may arise for that position even when the Cache is not full.

For eg, instructions program may start in block 1 and continue in block 129, after a branch. As this program is executed, both of these blocks are transferred to the block-1 position in the cache. Contention is resolved by allowing the new block to overwrite the currently resident block.

placement of a block in the Cache is determined from the memory address. The memory address can be divided into three fields as shown in fig.

206

The low-order 4 bits select one of 16 words in a block.

When a new block enters the cache, the 7-bit cache block field determines the cache position in which this block must be stored.

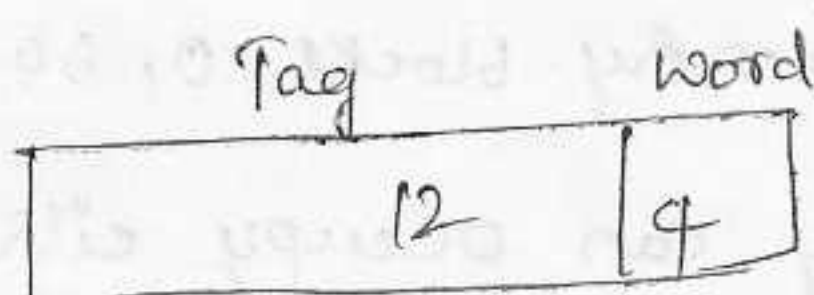
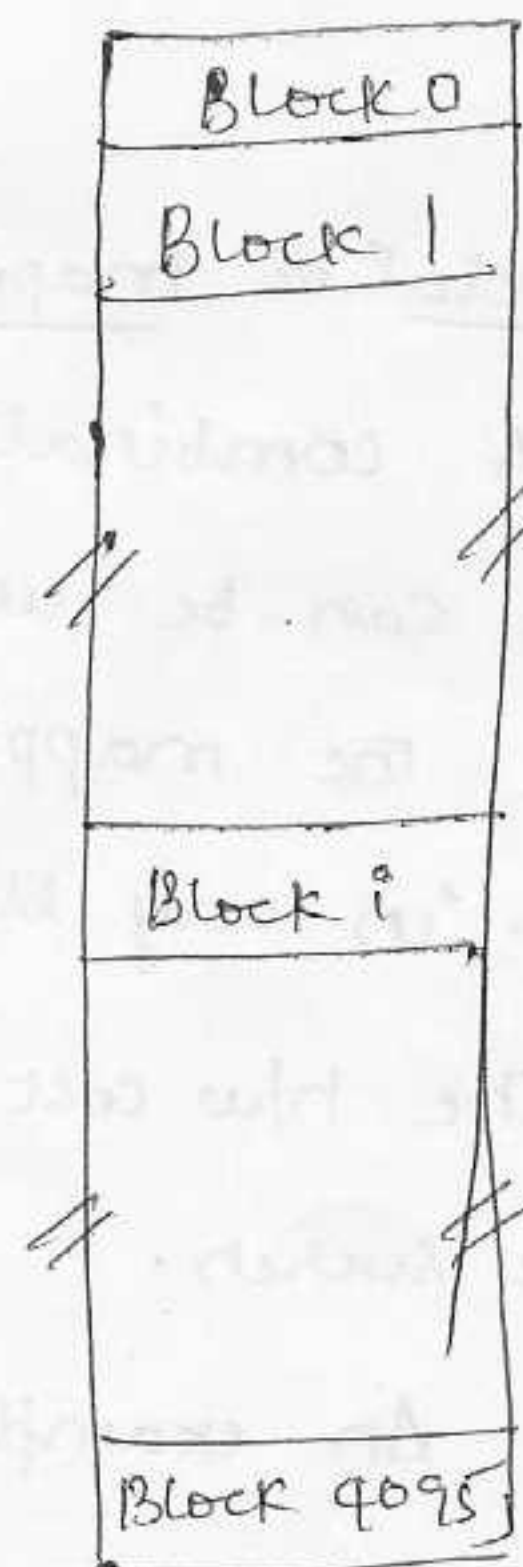
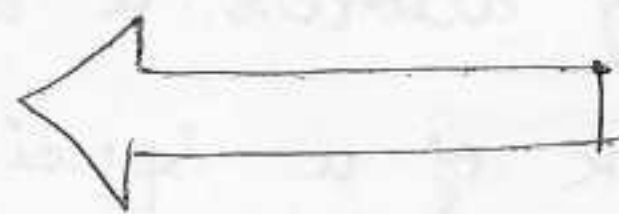
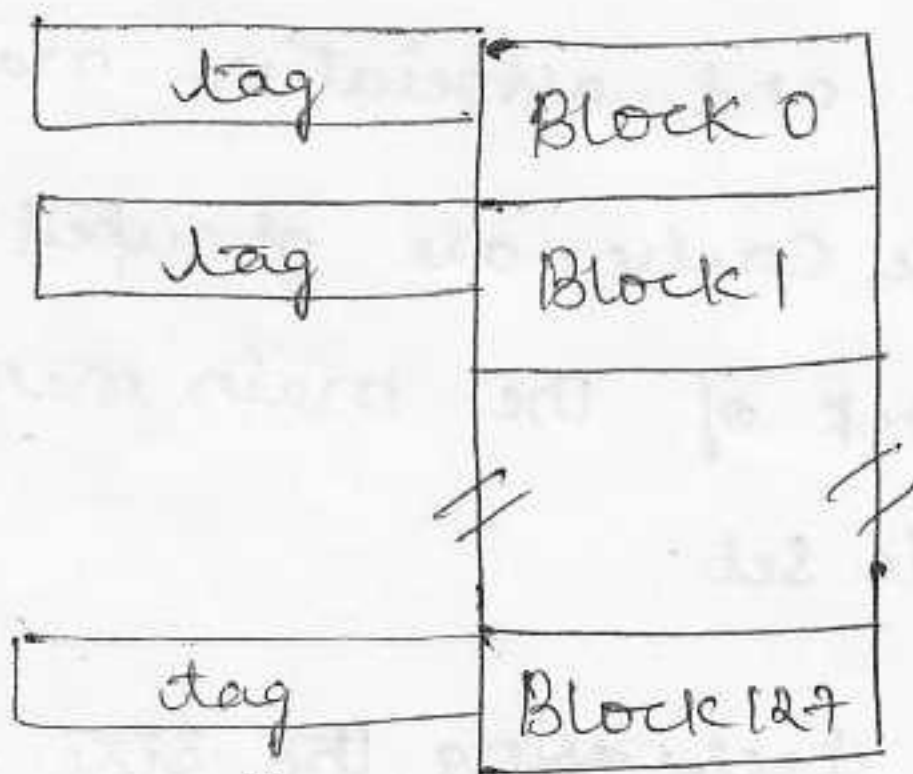
The high-order 5 tag bits associated with its location in the cache. They identify the 32 blocks that are mapped into cache.

As execution proceeds, the 7-bit cache block field of each address generated by the processor points to a particular block location in the cache. The high-order 5 bits of the address are compared with the tag bits associated with the cache location. If they match, then the desired word is in that block of cache.

If there is no match, then the block containing the required word must first be read from the main memory and loaded into the cache.

The direct-mapping is easy to implement but it is not very flexible.

Associative mapping :-



Main Memory address

This is a much more flexible mapping method, in which a main memory block can be placed into any cache block position.

In this 12 tag bits are required to identify a memory block when it is resident in the cache. The tag bits of an address received from the processor are compared to the tag bits of each block of the cache, if the desired block is present. This is called the 'associative-mapping' technique.

A new block that has to be brought into the cache has to replace an existing block only if the cache is full. For this case we need an algorithm to select that particular block to be replaced.

The cost of an associative cache is higher than the cost of a direct-mapped cache because of the need to search all 128 tag patterns to determine whether a given block is in the cache.

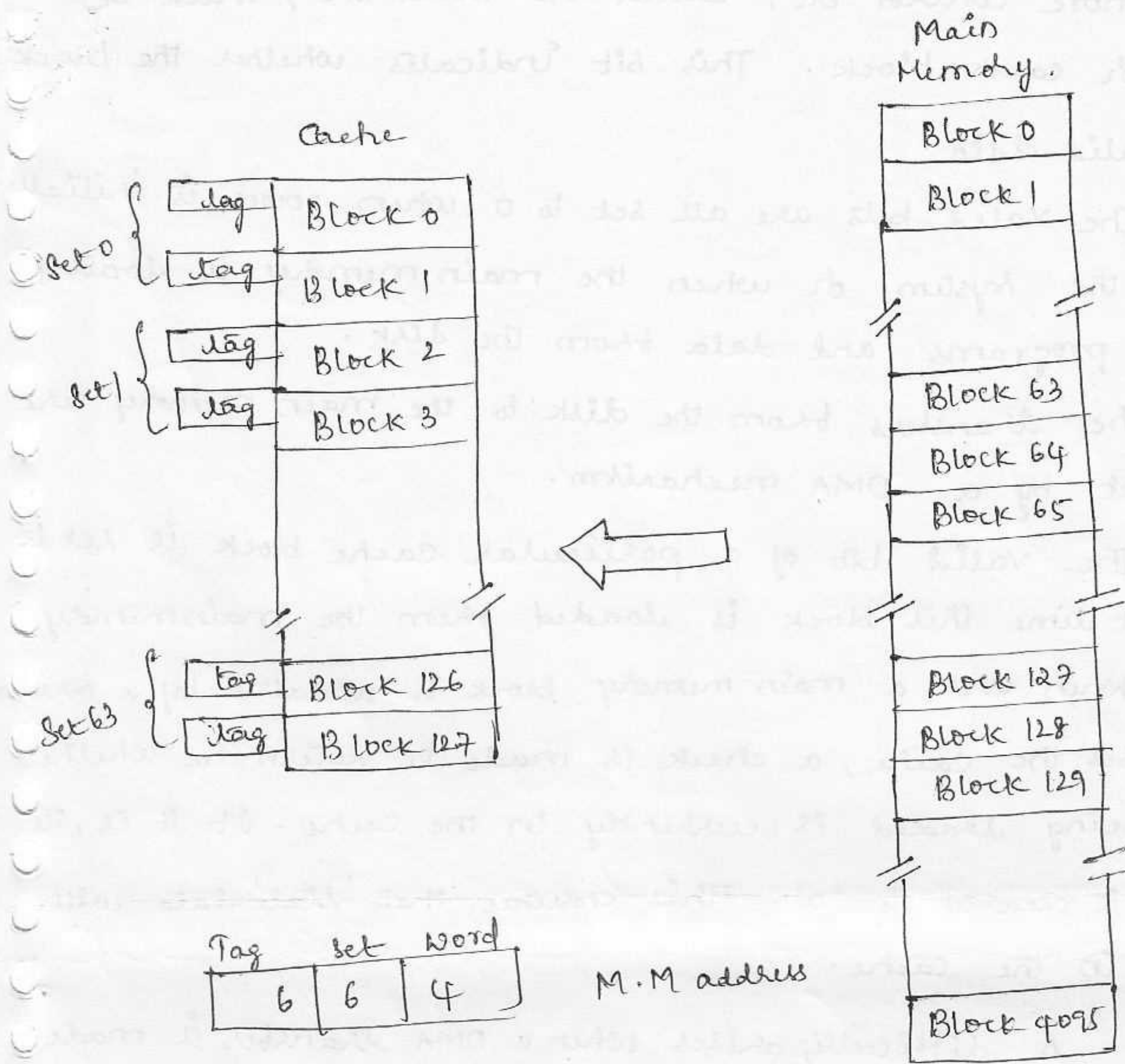
A search of this kind is called an associative search.

Set-Associative mapping :-

A combination of the direct and associative mapping techniques can be used. Blocks of the cache are grouped into sets and the mapping allows a block of the main memory to reside in any block of a specific set.

The h/w cost is reduced by decreasing the size of the associative search.

An example is shown in the fig. For a cache with two blocks per set. In this case memory blocks 0, 64, 128, ..., 4032 map into cache set 0 and they can occupy either of the two block positions within this set.



Having 64 sets means that the 6-bit set field of the address determines which set of the cache might contain the desired block. The tag field of the address must then be associatively compared to the tags of the two blocks of the set to check if the desired block is present.

For the main memory and cache sizes, four blocks per set can be accommodated by a 5-bit set field, eight blocks per set can be accommodated by a 4-bit set field and so on. The extreme condition of 128 blocks per set requires no set bits & corresponds to the fully associative technique, with 12 tag bits. The other extreme of one block per set is the direct mapping method. A cache that has K blocks per set is referred to as a K-way set-associative cache.

one more control bit, called the valid bit, must be provided for each block. This bit indicates whether the block contains valid data.

The valid bits are all set to 0 when power is initially applied to the system or when the main memory is loaded with new programs and data from the disk.

The transfers from the disk to the main memory are carried out by a DMA mechanism.

The valid bit of a particular cache block is set to '1' the first time this block is loaded from the main memory.

When ever a main memory block is updated by a source that bypasses the cache, a check is made to determine whether the block being loaded is currently in the cache. If it is, its valid bit is cleared to '0'. This ensures that 'stale' data will not exist in the cache.

A difficulty arises when a DMA transfer is made and the cache uses the write-back protocol.

One solution to this problem is to flush the cache by forcing the dirty data to be written back to the memory before the DMA transfer takes place.

The need to ensure that two different entities (processor & DMA) uses the same copies of data is referred to as cache-coherence problem.

Replacement algorithms :-

When a block is to be overwritten, it is sensible to overwrite the one that has gone the longest time without being referenced. This block is called the last recently used (LRU) block and the technique is called the LRU Replacement algorithm.

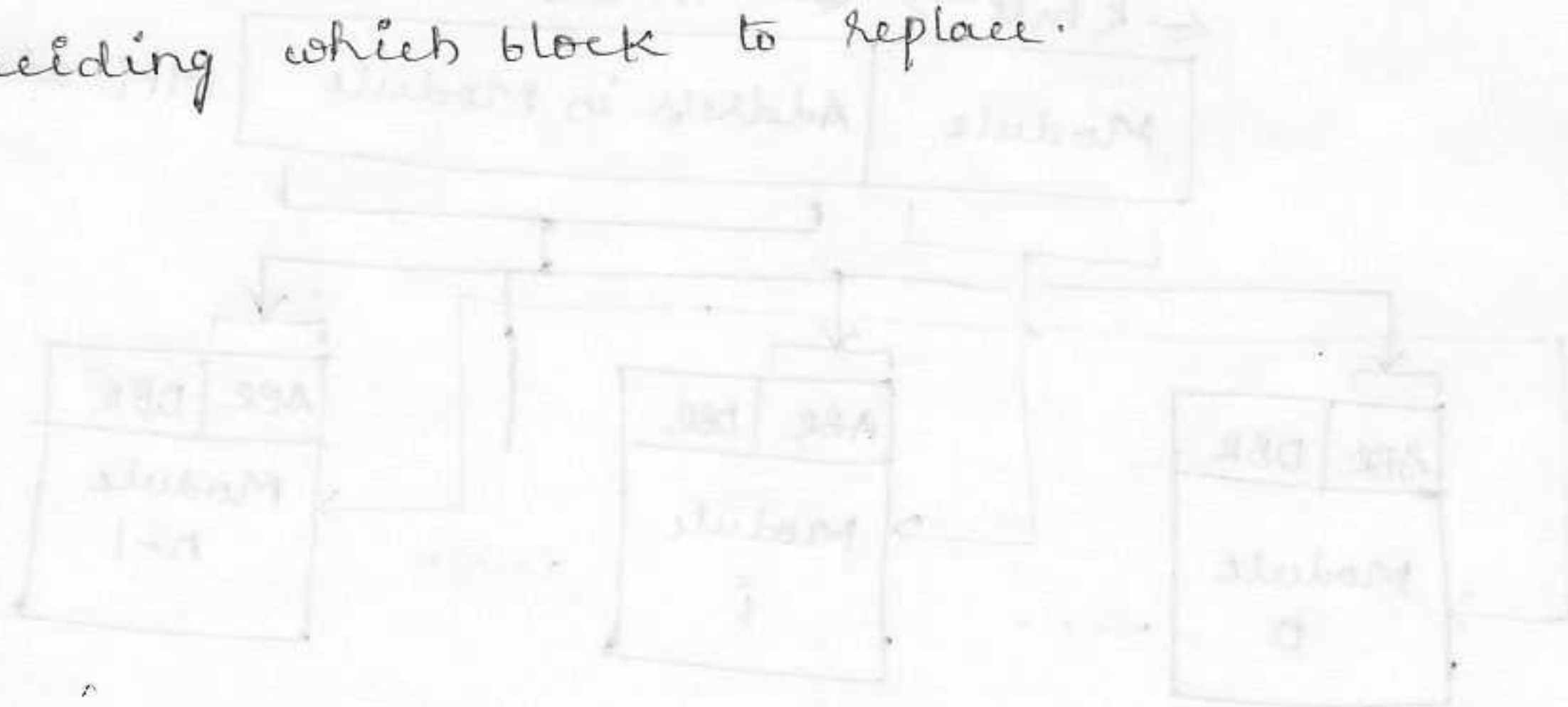
LRU Algorithm:-

To use the LRU algorithm, the cache controller must track references to all blocks. To track the LRU block of a four-block set in a set associative cache. A 2-bit counter can be used for each block. When a 'hit' occurs, the counter of the block that is referenced is set to 0.

When a 'miss' occurs and the set is not full, the counter associated with the new block loaded from the main memory is set to 0 and the values of all other counters are increased by one.

When a 'miss' occurs and the set is full, the block with the counter value 3 is removed, the new block is put in its place and its counter is set to '0'. The other three block counters are incremented by one.

The LRU algorithm has been used extensively. Although it performs well for many patterns, it can lead to poor performance in some cases. For eg, LRU produces bad results, when accesses are made to sequential elements of an array that is too large to fit into the cache. Performance of the LRU algorithm can be improved by introducing a small amount of randomness in deciding which block to replace.



performance Considerations :-

Two key factors of a computer are performance & cost.

performance depends on how fast machine instructions can be brought into the processor for execution & how fast they can be executed.

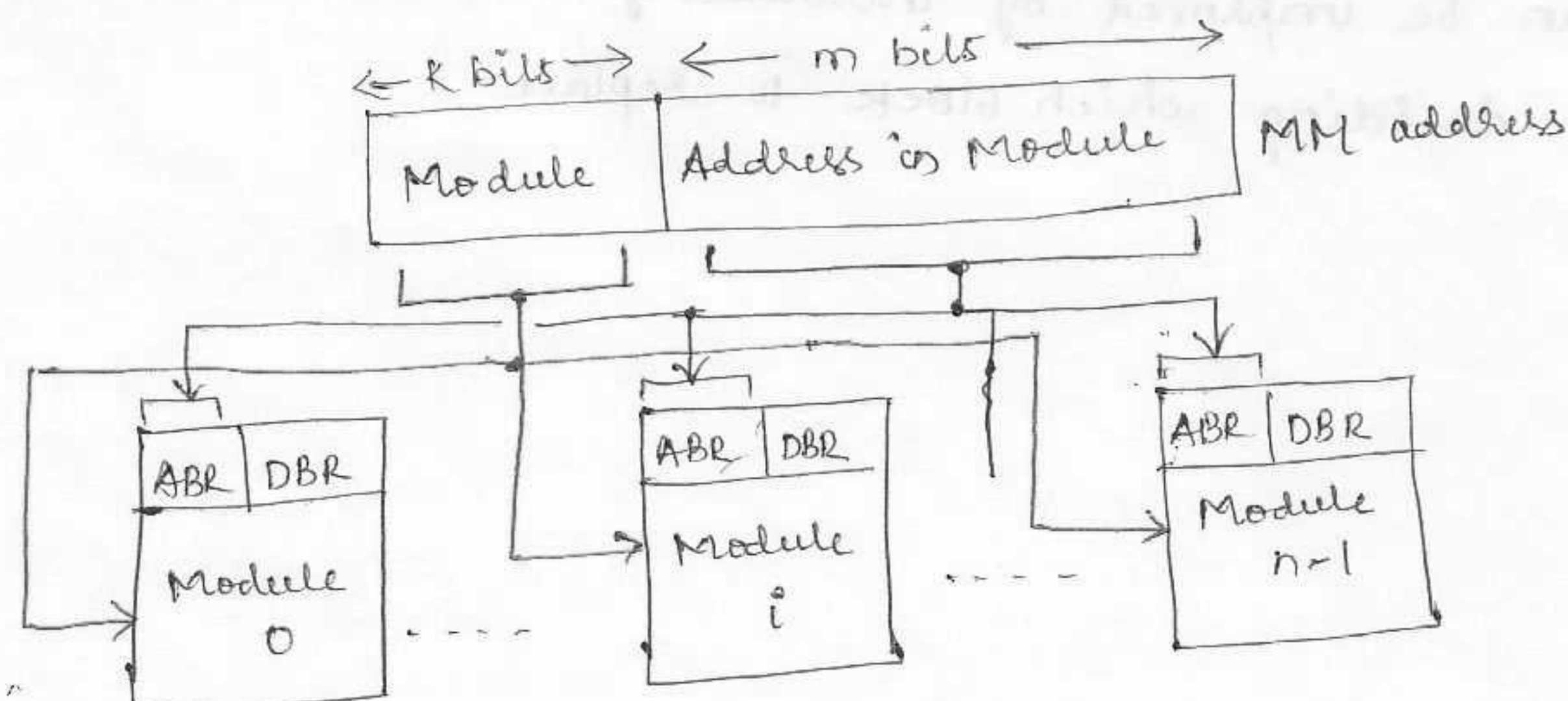
An effective way to introduce parallelism is to use an interleaved organization.

Interleaving :-

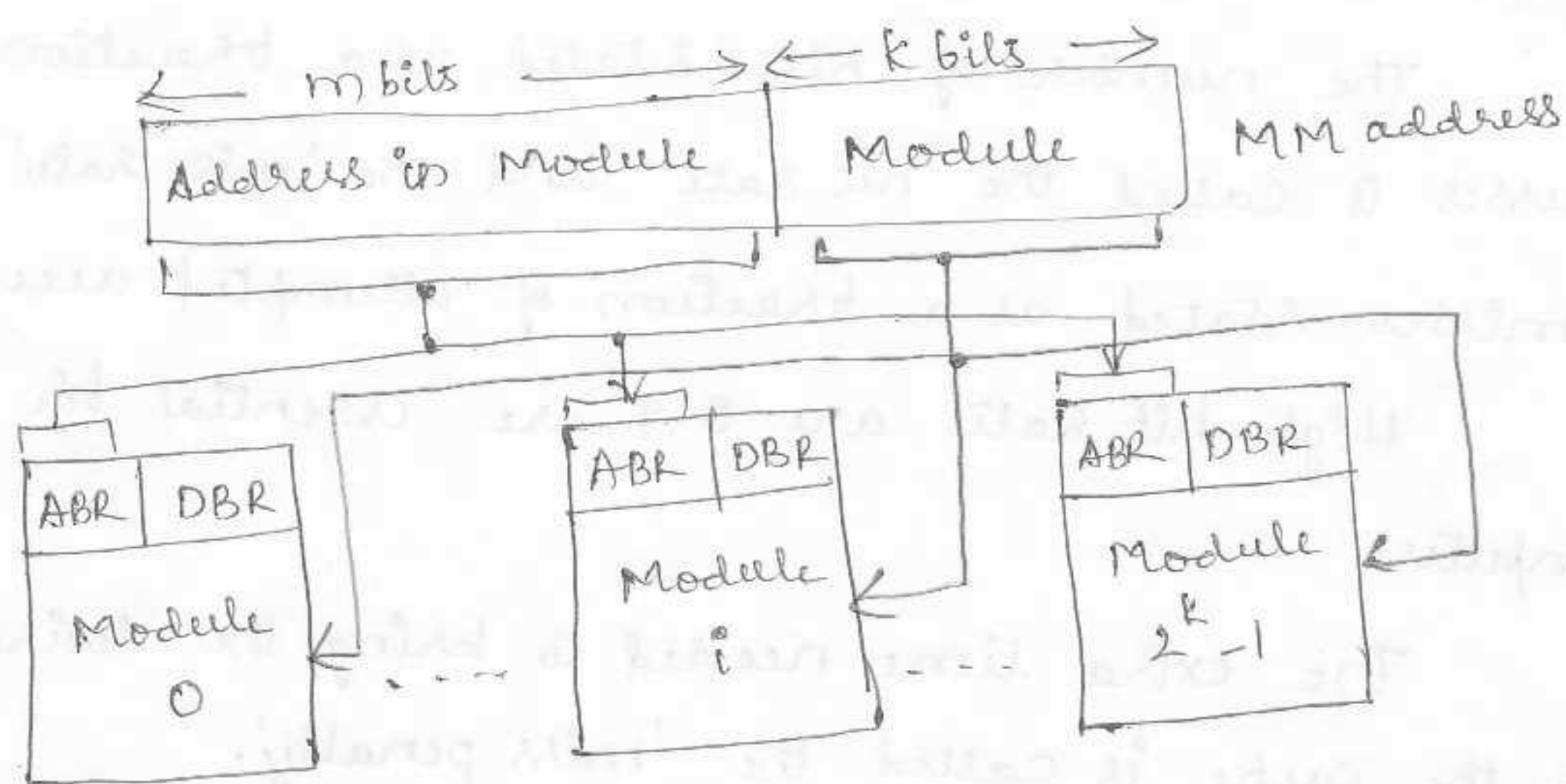
If the main memory of a computer is structured as collection of physically separate modules, each with its own address buffer register (ABR) and data buffer register (DBR),

The memory access operations may proceed in more than one module at the same time. Thus, the aggregate rate of transmission of words to and from the main memory system can be ⁱⁿ decreased.

Two methods of address layout are shown below.



a) Consecutive words in a module.



b) Consecutive words in consecutive modules

In fig (a), the memory address generated by the processor is decoded. The high-order k bits name one of ' n ' modules and the low-order ' m ' bits name a particular word in that module. When consecutive locations are accessed, when a block of data is transferred to a cache, only one module is involved. At the same time, however, devices with DMA ability may be accessing information in other memory modules.

The fig (b) is a more effective way to address the modules is called Memory interleaving. The low-order k bits of the memory address select a module, and the high-order m bits name a location within that module. The consecutive addresses are located in successive modules.

Any component of the system that generates requests for access to consecutive memory locations can keep several modules busy at any one time.

This results in both faster access to a block of data & higher avg. utilization of the memory system as a whole. To implement the interleaved structure, there must be 2^k modules. Otherwise, there will be gaps of nonexistent locations in the memory address space.

Hit rate & Miss penalty :-

The number of hits stated as a fraction of all attempted accesses is called the 'hit rate' and the 'miss rate' is the number of misses stated as a fraction of attempted accesses.

High hit rates are 0.9 are essential for high-performance computers.

The extra time needed to bring the desired information into the cache is called the 'miss penalty'.

In general, the miss penalty is the time needed to bring a block of data from a slower unit in the memory hierarchy to a faster unit.

The miss penalty is reduced if efficient mechanism for transferring data b/w the various units of the hierarchy are implemented.

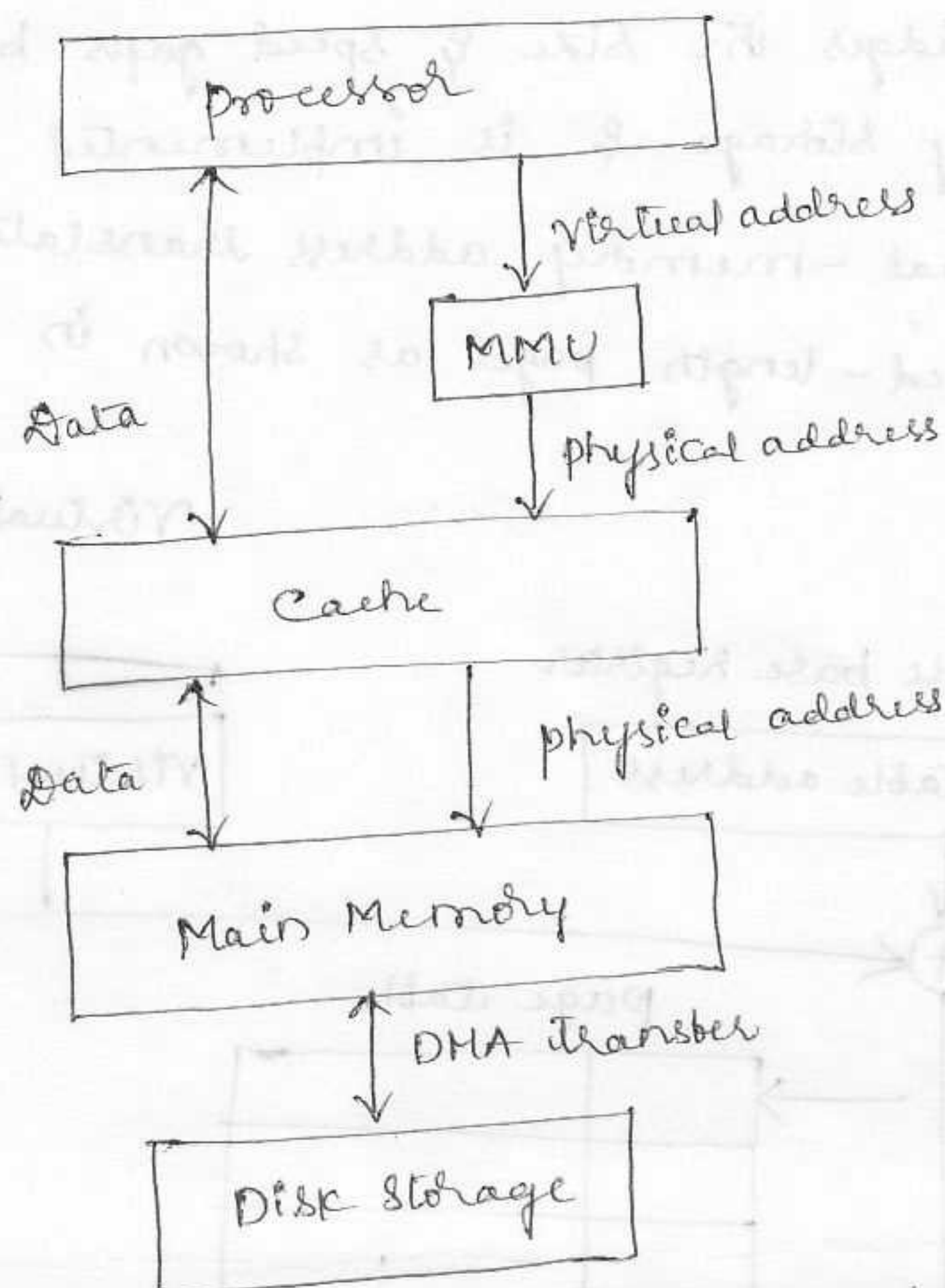
Virtual memories :-

Techniques that automatically move program and data blocks into the physical main memory, when they are required for execution is called virtual techniques.

The binary addresses that the processor issues for either instructions or data are called virtual or logical addresses. These addresses are translated into physical addresses by a combination of h/w & s/w components.

If a virtual address refers to a part of the program or data space that is currently in the physical memory, then the contents of the appropriate location in the main memory are accessed immediately. If the referenced address is not in the main memory, its contents must be brought into a suitable location in the memory.

The below fig. shows a organization that implements virtual memory. A special h/w unit called the memory management



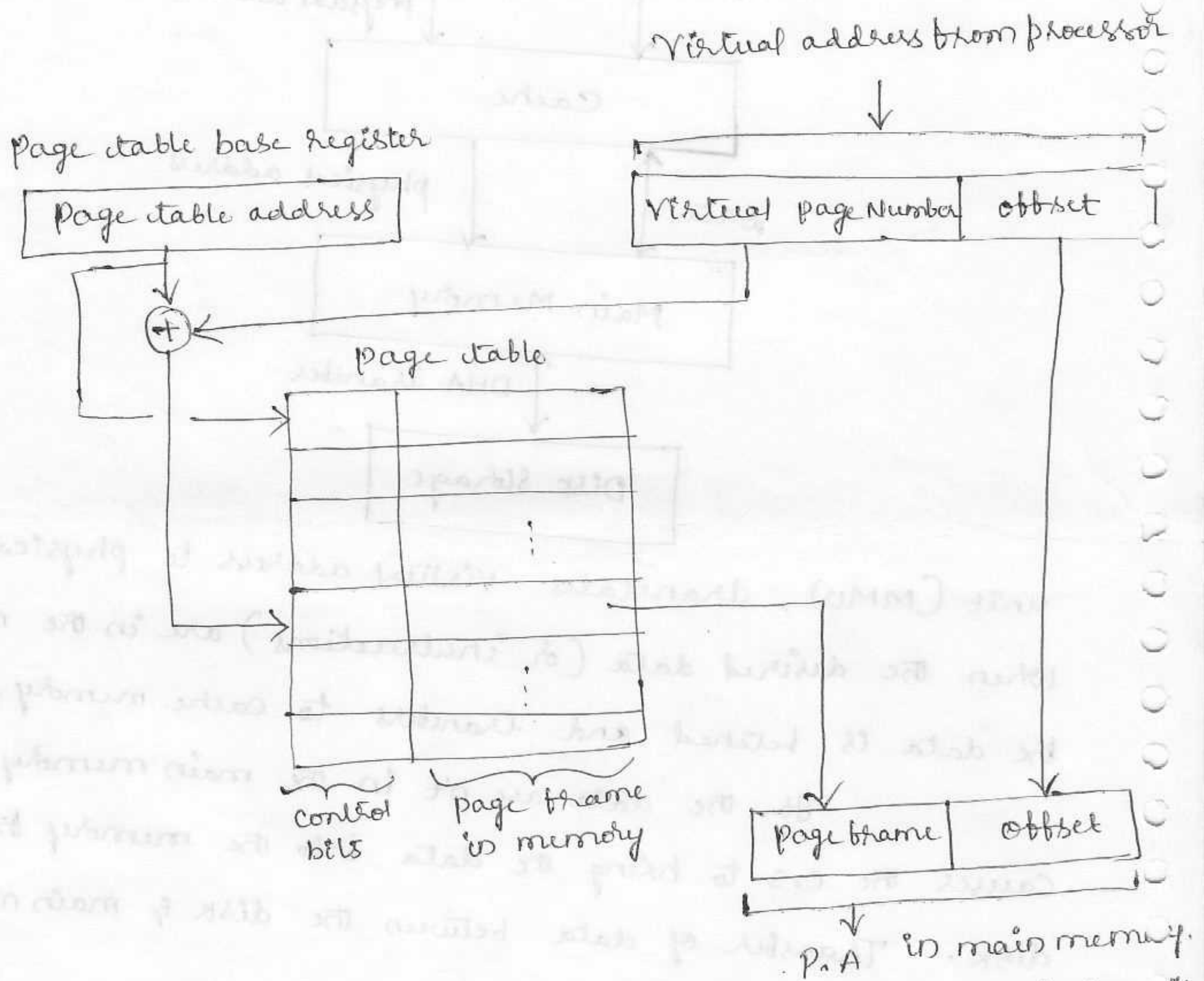
unit (MMU), translates virtual address to physical address. When the desired data (or instructions) are in the main memory the data is fetched and transfers to cache memory. If the data are not in the main memory, the MMU causes the O.S to bring the data into the memory from the disk. Transfer of data between the disk & main memory is performed using DMA.

Address Translation :-

A simple method for translating virtual addresses into physical addresses is to assume that all programs and data are composed of fixed-length units called pages, each of which consists of a block of words that occupy contiguous locations in the main memory.

The cache memory bridges the ^{speed} gap b/w the processor and the main memory and is implemented in h/w. The virtual-memory mechanism bridges the size & speed gaps b/w the main memory and secondary storage & is implemented in s/w.

A virtual-memory address translation is based on the concept of fixed-length pages as shown in fig.



Each virtual address generated by the processor, whether it is for an instruction fetch or an operand fetch/store operation is interpreted as a virtual page Number (High-order bits) followed by an offset (low-order bits) that specifies the location of a particular byte (or word) within a page.

This information that ~~indicates~~ an area in the main memory that can hold one page is called a 'page frame'.

The starting ^{address} of the page table is kept in a 'page table base reg'.

216
By adding the virtual page number to the contents of this register, the address of the corresponding entry in the page table is obtained.

Each entry in the page table also includes some control bits that describe the status of the page while it is in the main memory.

One bit indicates the validity of the page, whether the page is loaded in main memory.

Another bit indicates whether the page has been modified during resident in the memory.

Other control bits indicate various restrictions that may be imposed on accessing the page.

The page table information is used by the MMU for every read & write access, so the page table should be situated within the MMU.

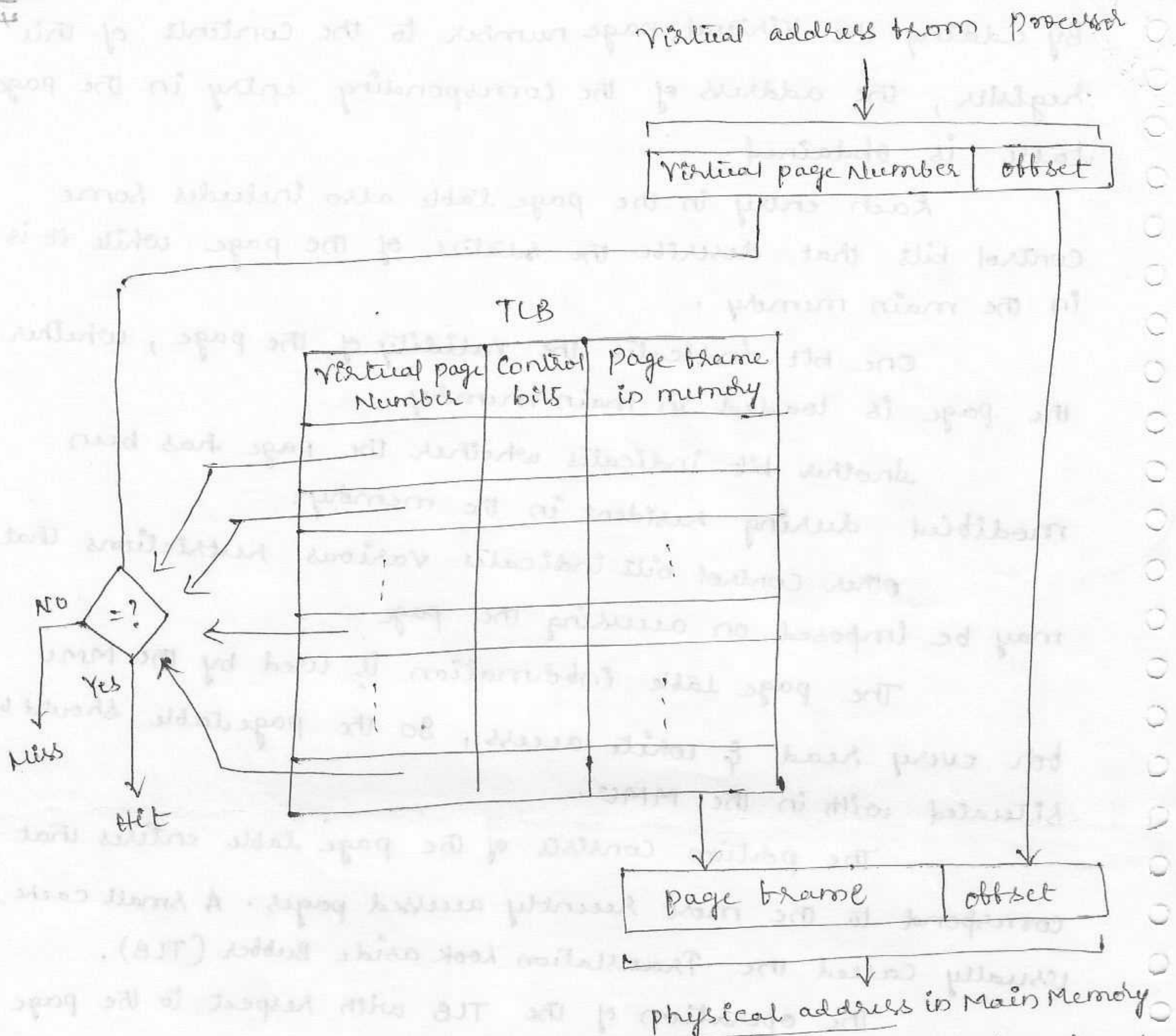
The portion consists of the page table entries that correspond to the most recently accessed pages. A small cache, usually called the Translation Lookaside Buffer (TLB).

The operation of the TLB with respect to the page table in the main memory is same as the cache memory. The information that constitutes a page table entry, the TLB must also include the virtual address of the entry.

The fig. shows the organization of a TLB where the associative mapping technique is used.

The contents of the TLB be coherent with the content of page tables in the memory. When the O.S. changes the contents of page tables, it must simultaneously invalidate the corresponding entries in the TLB.

One of the control bits in the TLB is provided for this purpose. When an entry is invalidated, the TLB will



acquire the new information as a part of the MMU's normal response to access misses.

Address translation may follow in this way

- 1) Given a virtual address, the MMU looks for this referenced page in the TLB.
- 2) If the page table entry for this page is found in the TLB, the physical address is obtained immediately.
- 3) If there is a miss in the TLB, then the required entry is obtained from the page table in the main memory and the TLB is updated.

When a program generates an access request to a page, if that page isn't in the main memory, then a 'page fault' is occurred.

A page fault can be detected, by the MMU asks the O.S. to intervene by raising an exception (interrupt). processing of the active task is interrupted, the control is transferred to the O.S.

The O.S. then copies the requested page from the disk into the main memory and returns control to the interrupted task.

A long delay occurs while the page transfer takes place.

If a new page is brought from the disk when the main memory is full, it must replace one of the resident pages.

The problem of choosing which page to remove is critical, because main memory has larger memory. We should large portions in main memory. This will reduce the frequency of transfers to & from the disk.

Here also, the LRU replacement algorithm is used. and the control bits in the page table entries can indicate usage.

One simple scheme is, the control bit is set to '1' when ever the corresponding pages is referenced. The O.S, clears this bit in all page-table entries, thus providing a simple way of determining which have not been used recently.

A modified page has to be written back to the disk before it is removed from the main memory.

The address translation process in the MMU requires some time to perform, mostly dependent on the time needed to look up entries in the TLB.

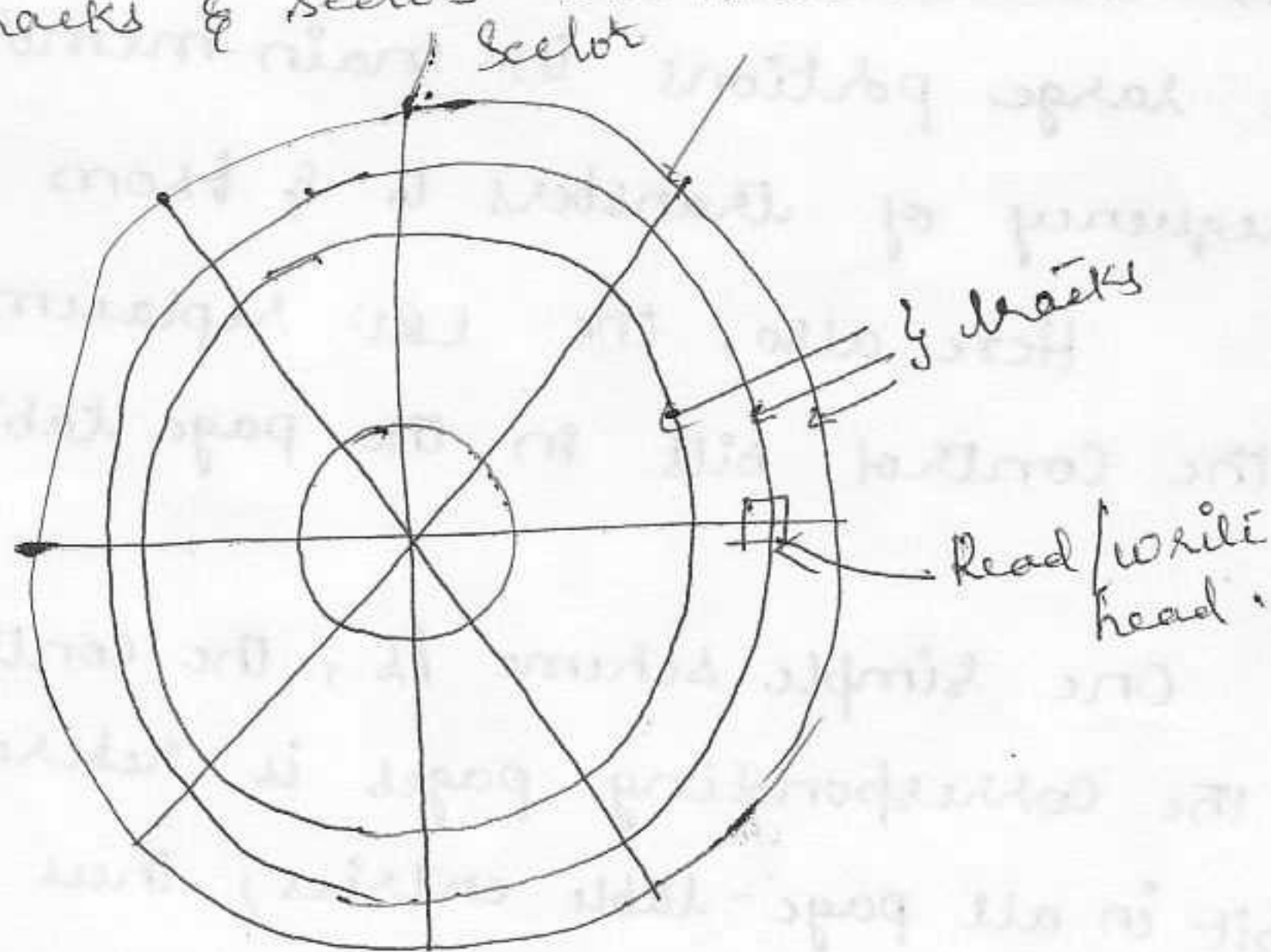
219

We can reduce the average translation time by including one or more special registers that retain the virtual page number and the physical page frame of the most recently performed translations. The information in these registers can be accessed more quickly than the TLB.

Magnetic Disks :-

— Magnetic disk is a circular plate constructed of metal or plastic coated with magnetized material. The both sides of the disk are used & several disks may be stacked on one spindle with read/write heads available on each surface.

All disks rotate together at high speed and are not stopped. Bits are stored in the magnetized surface in spots along concentric circles called tracks. The tracks are commonly divided into sections called sectors. The min. quantity of information which can be transferred is a sector. The subdivision of one disk surface into tracks & sectors are shown in below fig.



Some units use a single read/write head for each disk surface. The track address bits are used by mechanical assembly to move the head into the specified track position before read/write. In other disk systems, separate read/write heads are provided for each track in each surface. The address bits can then select a particular track electronically thro' a decoder circuit.

This type of unit is more expensive & is found only in very large computer systems.

permanent timing tracks are used in disks to synchronize the bits & recognize the sectors. A disk system is addressed by address bits that specify the disk number, disk surface, the sector number and track within the sector.

Information transfer is very fast. Disks may have multiple heads & simultaneous transfer of bits from several tracks at the same time.

Disks that are permanently attached to unit & can't be removed by the user are called hard disks. A disk drive with removable disks is called a floppy disk.

Magnetic tape :-

A magnetic tape transport consists of the electrical, mechanical and electronic components to provide the parts & control mechanism for a magnetic tape unit.

The tape itself is a strip of plastic coated with a magnetic recording. Bits are recorded as magnetic spots on the tape along several tracks. Seven or nine bits are recorded simultaneously to form a character together with a parity bit. Read/write heads are mounted one in each track so that data can be recorded & read as a sequence of characters.

Magnetic tape units can be ~~started~~ stopped, started to move forward or in reverse. The information is recorded in blocks referred to as records. Gaps of unrecorded tape are inserted b/w records where the tape can be stopped.

The tape starts moving while in a gap and attains its constant speed.

By reading the bit pattern at the end of the record, the control recognizes the beginning of a gap.

A tape unit is addressed by specifying the record numbers & no. of records. Records may be fixed or variable length.

RAID :-

RAID (Redundant Array of ^{expensive} ~~Independent~~ Disks) is a Technology where multiple independent disks are maintained. In this, a single large file is partitioned & each partition is placed in each of these independent disks. Whenever any process requires these partitions, it can easily access these disks parallelly & acquire the required data.

When these disks are accessing parallelly, hence the given accessing time is decreased where as bandwidth of data transfer is increased.

The data belonging to a single large file, when gets partitioned then, each partition is referred as 'stripe' ~~and~~ used in way each of these strips are placed in these disks is referred as striping.

Hence RAID is dependent on two factors.

- 1) Striping of data in order to access all the disks parallelly &
- 2) Data redundancy to increase reliability.

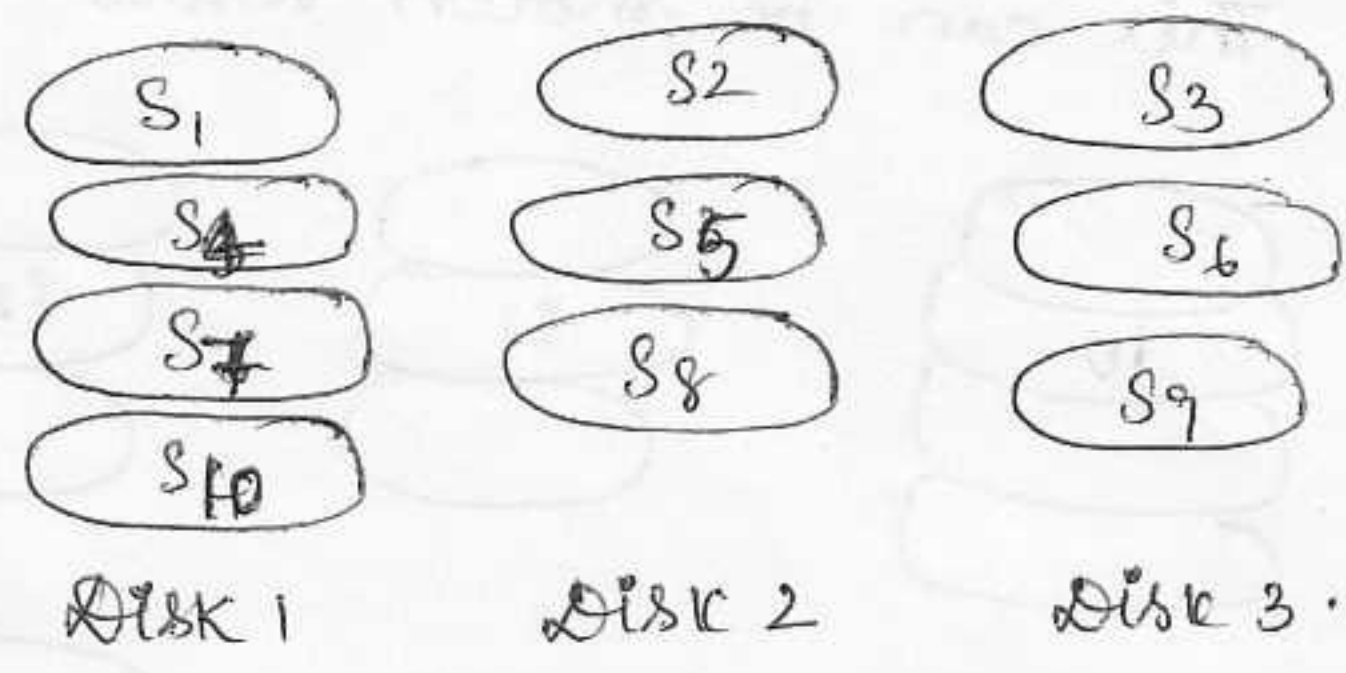
RAID has different levels which are explained below.

RAID Level 0 :-

In this level, stripes of data are belonging to the files is introduced in each disks, such that first strip is added in first disk, second strip is added in disk 2 and so on. As the no. of disks ends, the same procedure is repeated by considering disk 1 as initiating disk.

The stripping procedure ends only when all strips are filled in these disks. In Raid level 0, there is no provision for data redundancy.

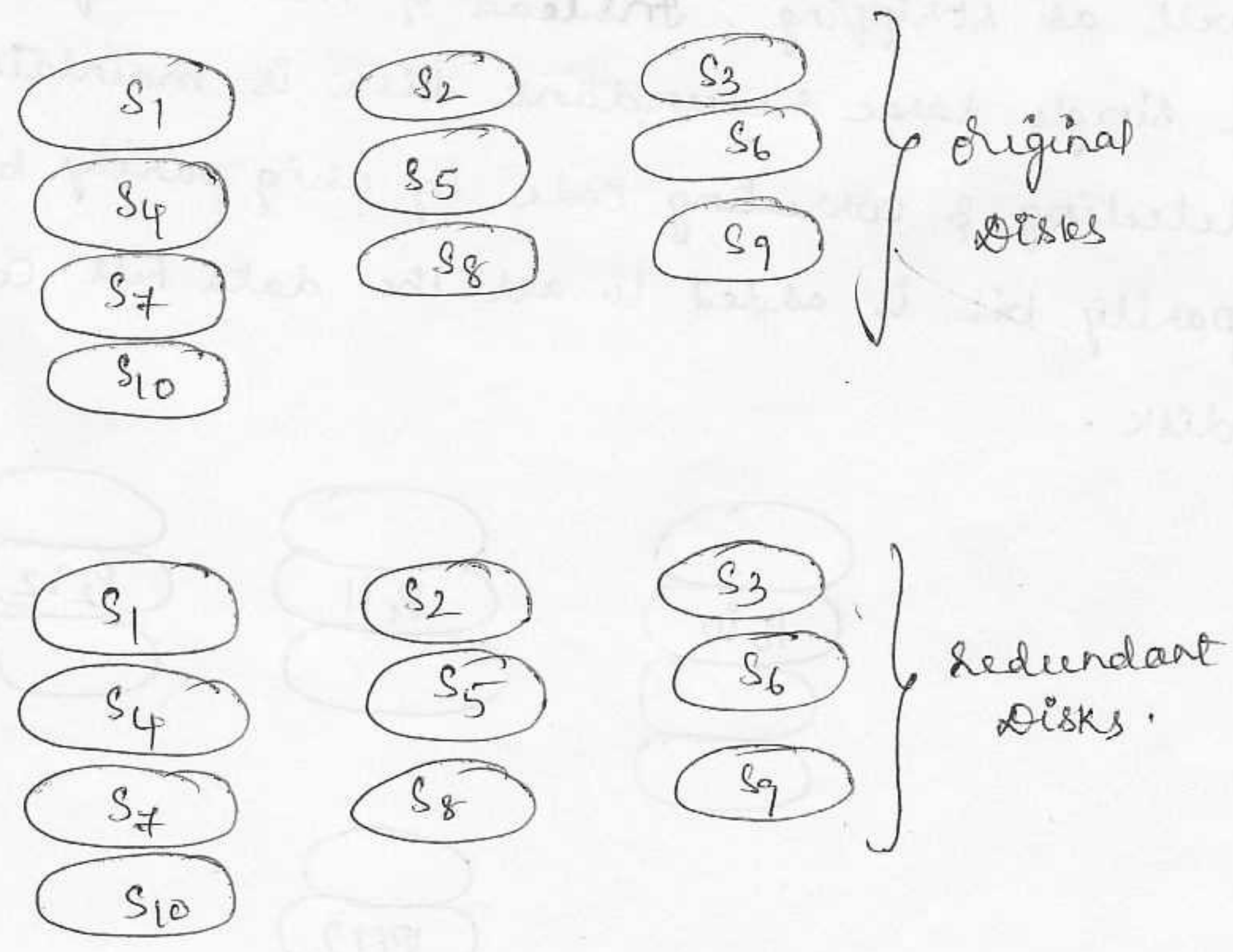
Various consequences occurring at RAID level 0 can be represented below;



RAID level 1:-

At this level data redundancy is considered to be of primary focus. RAID level 0 is applied here, in order to achieve redundancy each physical disk is maintained twice. As a result each stripe of data will have its copy in a separate identical disk. Hence reliability can be achieved.

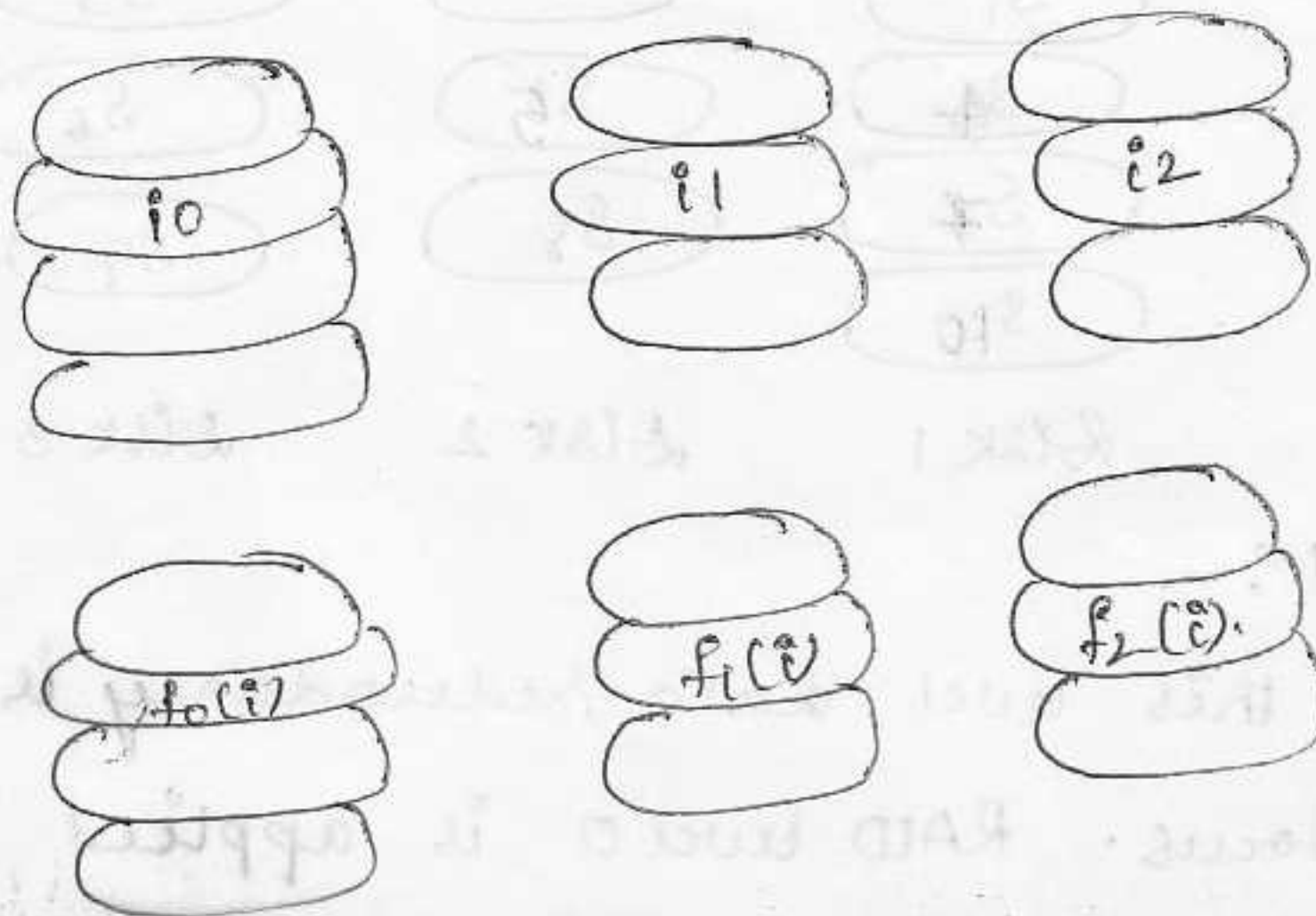
This can be shown below.



RAID level 2:

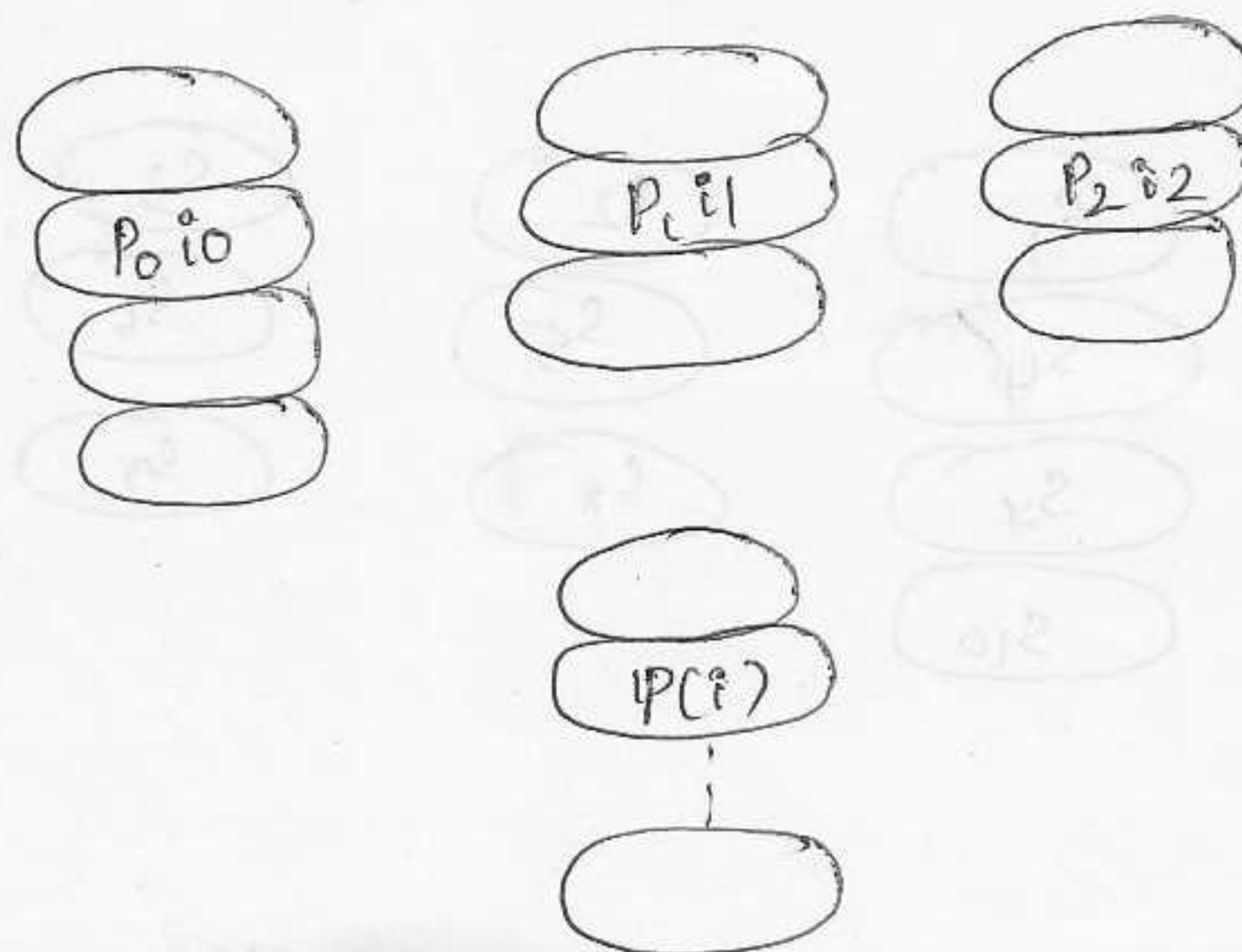
RAID level 2 is a combination of level 0 & level 1, Instead of looking the frequently occurring errors, an error detecting as well as correcting code is maintained across all the disks. The frequently used error detecting code is "Hamming code".

This can be shown below



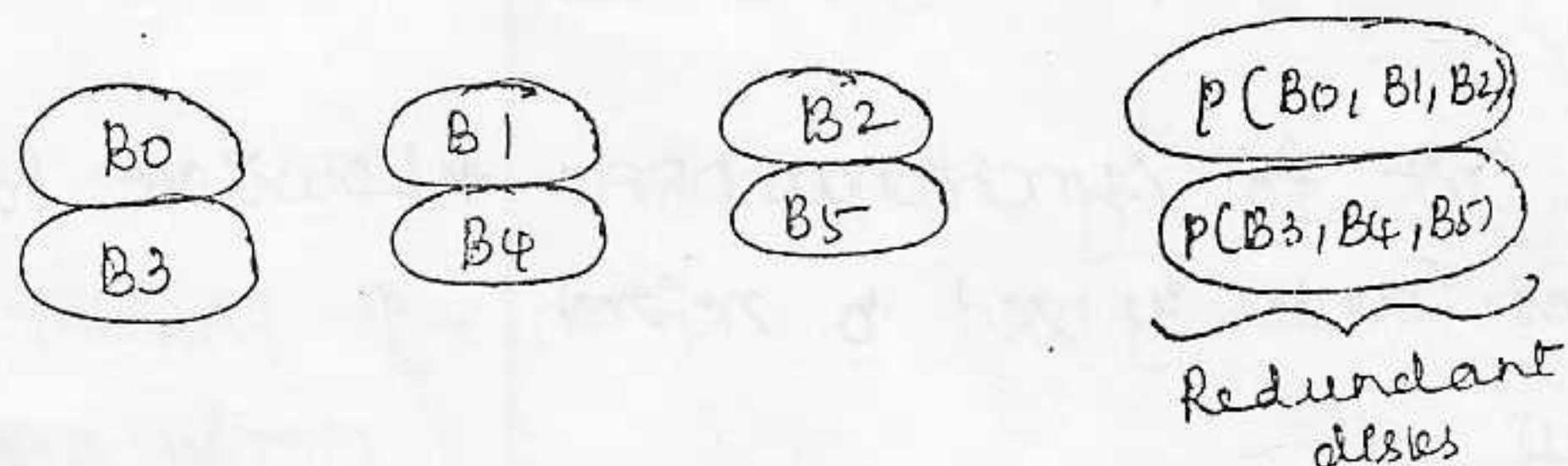
RAID level 3:

RAID level 3 reflects other levels in terms of redundancy as well as striping. Instead of maintaining equal no. of redundant a single large redundant disk is maintained. Also, the error detecting & correcting code by using parity bit concept (extra bit) parity bit is added to all the data bits corresponding to each disk.



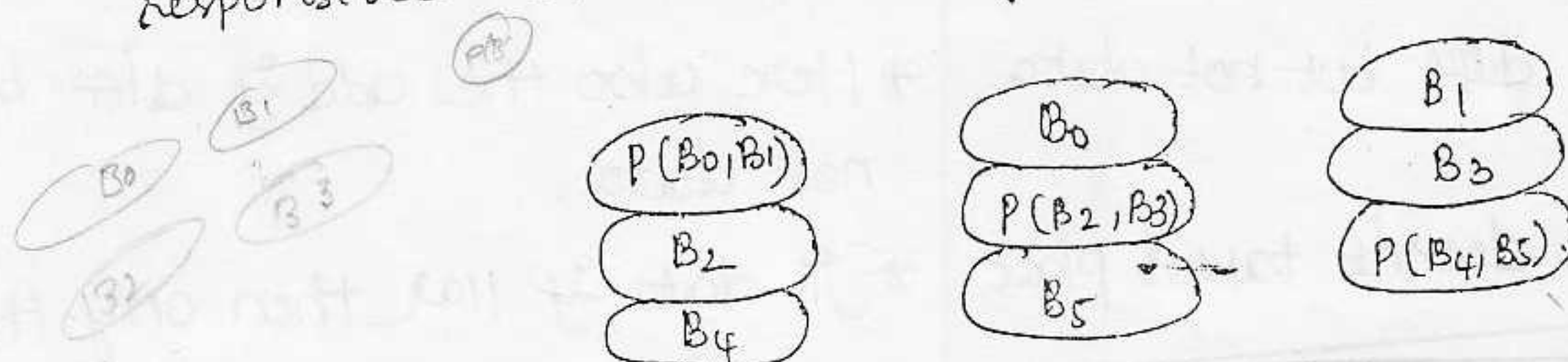
RAID level 4 :

Most of the data strips are combined to form a set of data. A single redundant disk is maintained. For each combination of blocks, a parity bit is added and the code placed in the redundant disk. In this level, the redundant disk with parity added data blocks turn out to be a less



RAID level 5 :

In order to overcome the limitation of level 4 i.e., instead of placing all the parity enabled blocks in one disk, they are distributed to all the disks. Hence all the disks now responsible for maintaining parity enabled blocks.



SRAM

- 1) Complex CKT
- 2) Costlier
- 3) Use latch & translational f/f's
- 4) Can be used as Cache
- 5) Temp. storage permanent
- 6) Current flow when it is accessed only
- 7) Hasn't latch
- 8) Less power consumption
- 9) we store automatically in latch

DRAM

- 1) Simple CKT
- 2) Cost effective
- 3) Capacitors used
- 4) Can be used as MM.
- 5) Temp storage
- 6) Charging Capacitors.
- 7) Hasn't latch
- 8) High power consumption.
- 9) Store data in auxiliary memory after discharging

Synchronous DRAM

- * In this it is directly synchronised with clock signal
- * The Complexity of CKT is not reduced
- * High Cost for synchronous DRAM
- * Refresh Counter is used to refresh the cell
- * Low density
- * Time delay is present
- * No data lines are used
- * Over write doesn't occur
- * Here add is diff but not data
- * If over write doesn't takes place
- * Here data i/p & data o/p reg's are present.

* It can be used with clock above 100MHz

SRAM

- 10) Manipulated by selection of bits
- 11) Continuous power supply is req'd
- 12) Fastest manipulation speed

Asynchronous DRAM

- * In this it is controlled asynchronously by using timing signal
- * The CKT Complexity is here
- * Low cost for asynchronous DRAM
- * It provides flexibility in designing memory systems.
- * High density
- * Time delay is not here
- * Data lines D₀ to D₇ are used
- * Over write occurs
- * Here also the add is diff but not data.
- * If data is 1101 then only the over write of data takes place
- * Here data i/p & data o/p reg's are not present.

DRAM

- 10) Manipulated by giving power supply.
- 11) Charging is req'd to capacitors
- 12) Slower manipulation.