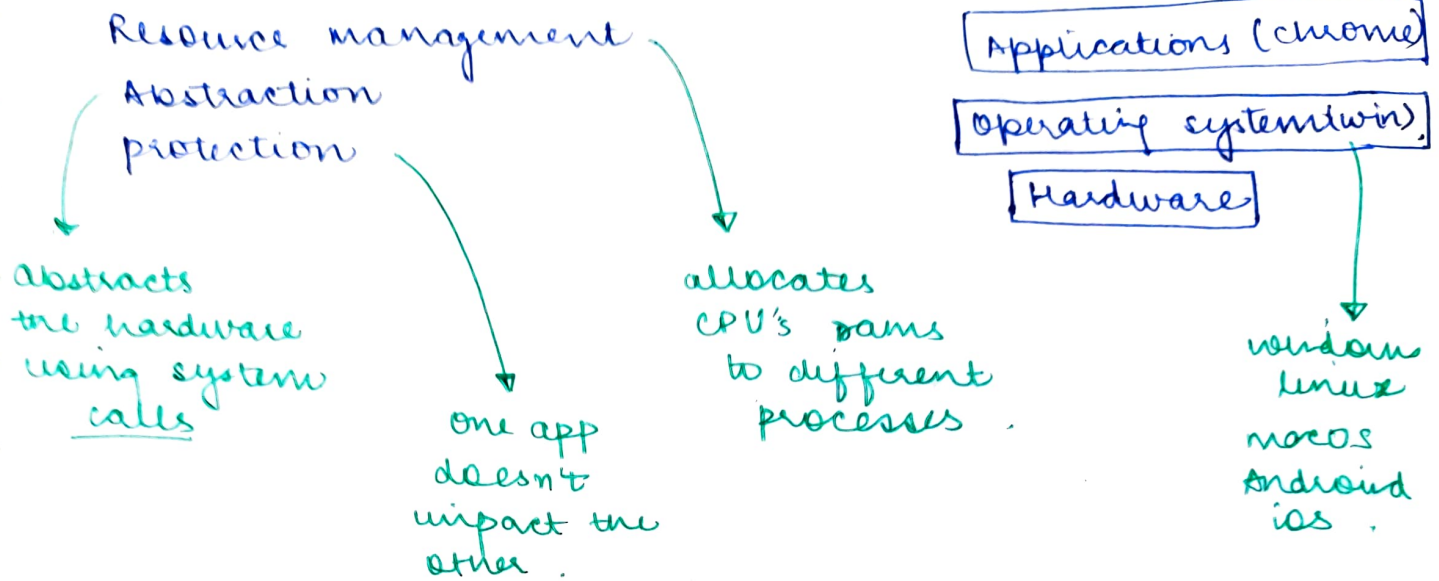


Operating System

- Applications have to interact with operating system to get things done.

services :-



Types of operating system

1) single tasking → MS-DOS (Inefficient)

└ one process is there in a RAM & running

2) Multiprogramming & Multitasking

- └ during computations process is assigned to the CPU
 - └ during I/O it's unassigned & assign some other process.
- multitasking is an extension of multiprogramming where process runs in time slots.

3) Multithreading

- └ using different threads for different processes
- ↳ smallest unit of execution.

4) Multiprocessing

- └ multiple processors in the system available

Multithreading

→ Deconvoluting something in browser and browser

Multiple things within a process.

multitasking

→ allowing to move & browsing web

two different processes running

Examples :-

word processors → saving / typing / formatting, spell checking etc.
IDE's, games etc

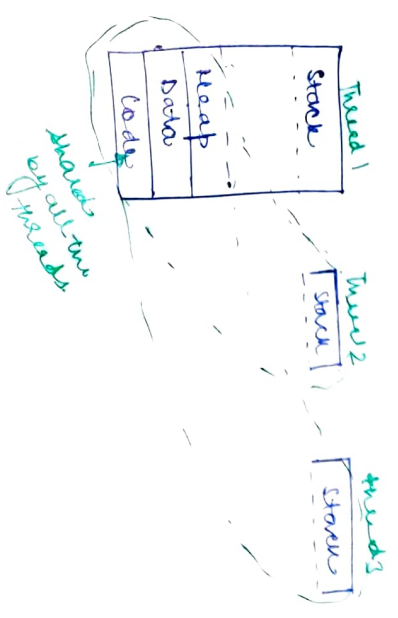
Disadvantages of multithreading

- Difficulty in writing / testing.
- Deadlock and race condition.
- Same resource is used by two threads.



Threads v/s Processes

Process → code, data, files, heap, stack etc
Thread → stack



Threads are

- faster to create / terminate.
- Multiple threads in a process
 - share address space
 - easier to communicate
 - context switching is equal.
- Threads are lightweight

→ consume less resources.

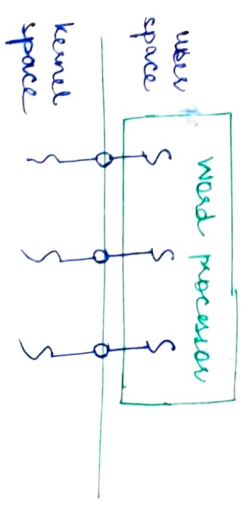
User Threads v/s Kernel Threads

- In user space
- Context switching faster
- one thread might block all other thread.
- cannot take advantage of multicore systems.
- creation is fast

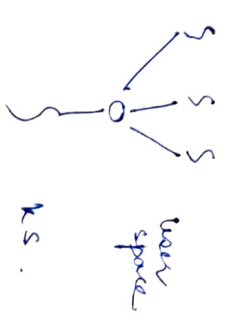
- in kernel space
- context switching slow
- A thread blocks itself only
- Take full advantage of multicore systems.
- slow.

Mapping of user threads to kernel threads.

one to one



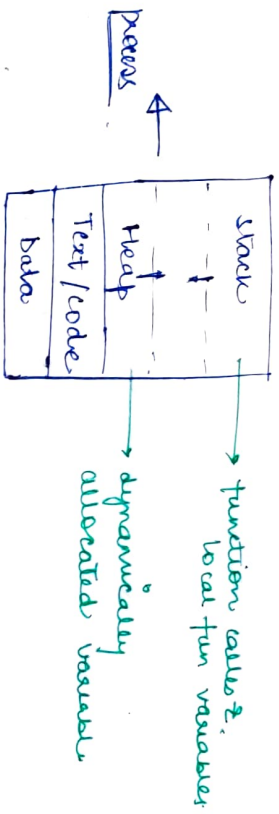
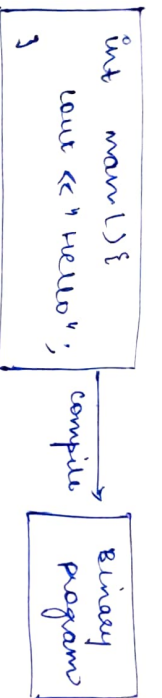
many to one



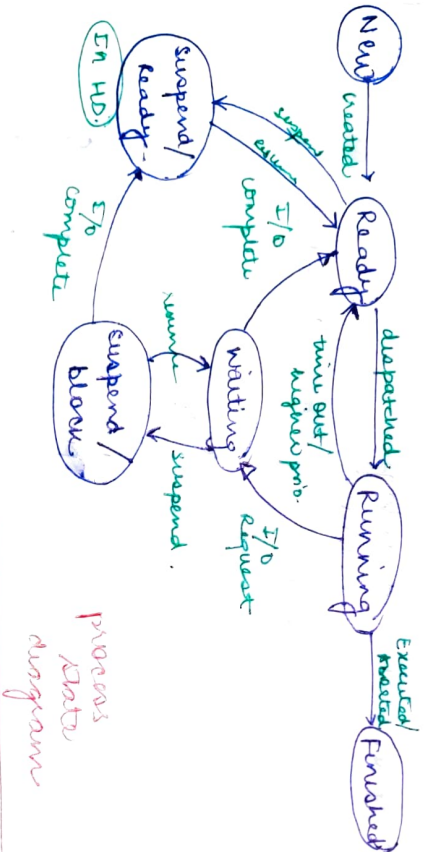
many to many



Program → set of instructions to perform a task



Process states



Process control block

- OS has to store all the variables and state of process before stopping it

For this we use PCB process control block.

- 1) process ID
- 2) process state
- 3) CPU registers
- 4) accounts information
- 5) I/O information
- 6) CPU scheduling information
- 7) memory information

Process scheduler

- 1) long term scheduler → which brings the process from disk to RAM.
- 2) short term scheduler → to ready state
- 3) medium term scheduler

The additional two states were managed by this suspend block & suspend ready.

↓

moves the process from ready state to running state (dispatcher)

Background of scheduling algorithm (short term scheduler)

→ different queues



→ short term scheduler and dispatcher pick one of the process from process ready queue

↓

des context switch and maintain memory state

→ when is a process picked by the scheduler.

(a) when some other process moves from running to waiting

(b) when some other process moves from running to ready

(c) when a new/extra existing process move to ready

(d) when process terminates.

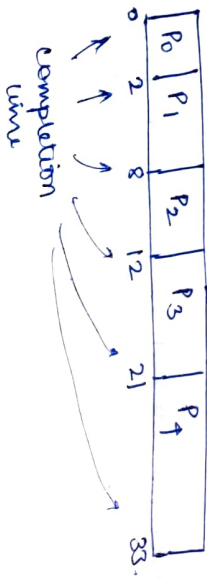
- arrival time
- completion time
- burst time
- time of CPU that process runs
- Turn around time
- waiting time → time spent in ready queue
- Response time → diff. b/w arrival time & first time it get CPU

Expectations from scheduling algorithm

- Max CPU utilization
- Max throughput
- Min turnaround time
- Min waiting time
- Min response time
- Fair CPU allocation

FCFS scheduling algorithm (Non-preemptive)

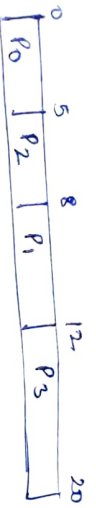
process	Arrival time	Execution time
P ₀	0	2
P ₁	1	6
P ₂	2	4
P ₃	3	9
P ₄	4	12



convoy effect → one CPU bound process
 2 many I/O bound processes
 ↓
 create big waiting queues

Shortest Job First (SJFS)

process	Arrival time	Burst time	comp. time	TAT	waiting
P ₀	0	5	5	5	0
P ₁	1	4	12	11	7
P ₂	2	3	8	6	3
P ₃	3	8	20	17	9



Non-preemptive

Preemptive Shortest Job First OR

Shortest Remaining time first

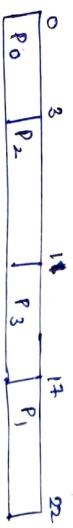
process	Arrival time	Burst time
P ₀	0	8
P ₁	1	7
P ₂	2	9
P ₃	3	5



* minimum average waiting time minimum
 + may cause high waiting and response time for CPU bound jobs.

Priority Scheduling

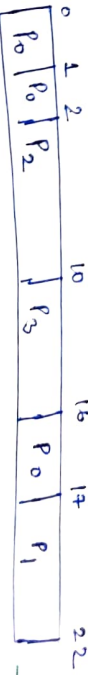
process	Arrival	priority	Burst
P ₀	0	5	3
P ₁	1	3	5
P ₂	2	15 (H)	8
P ₃	3	12	6



(Non-preemptive)

PREEMPTIVE

process	Arrival	priority	Burst
P ₀	0	5	8
P ₁	1	3	5
P ₂	2	15	8
P ₃	3	12	6



preemptive

Starvation of low priority

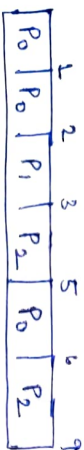
→ ageing → increase the priority with their age

Round Robin scheduling, (Preemptive)

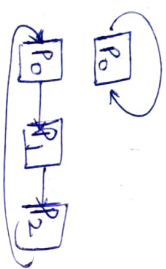
- Time quantum
- circular queue (Ready queue)
- Avg waiting time can be higher but good response time

process	Arrival	Burst
P ₀	0	3
P ₁	1	1
P ₂	1	5

time quantum = 2



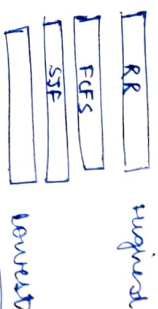
small context switch overheads
sensitive to time quantum
longer becomes FCFS



Multilevel queue scheduling

have different queues and apply different scheduling algorithms on them

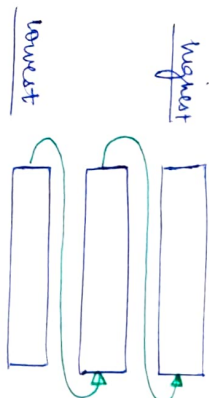
we can have priority queues where CPU is divided among queues on the basis of priority



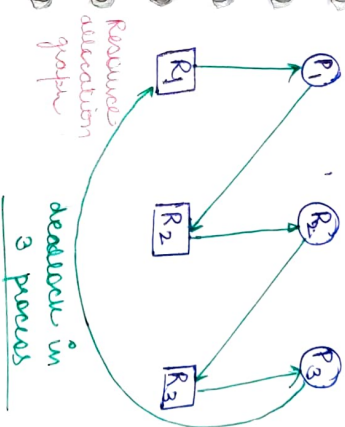
→ Starvation

Multilevel queue with feedback

Process should be allowed to switch between queues



DEADLOCK → happen when resources are non-shareable.

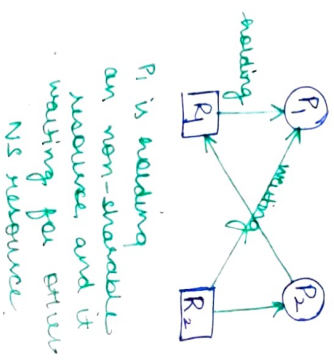
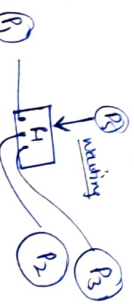


deadlock in 3 processes

necessary conditions

- 1) mutual exclusion → resources are non shareable
- 2) hold and wait → processes must be in hold & wait i.e. holding the resource & waiting for other resources
- 3) circular wait → processes are waiting in circular manner
- 4) NO preemption → resources can't be taken back by OS in the middle of process

there might be multiple instances of a resource



Deadlock Prevention methods

1) Deadlock prevention

Prevent one of the four necessary conditions for deadlock by assigning some rules. OS grant as the request but request are sent in a way that they never cause deadlock.

2) Deadlock avoidance

OS has to decide whether a request can be granted or not.

Ex - banker's Algorithm

3) Detection and Recovery

OS periodically runs a job and check if it has happened, if it has it is recovered by killing the process.

4) Ignore the detection ↳ simply ignore the deadlock.

Deadlock Prevention

prevent / eliminate one of the following four

- 1) Mutual Exclusion
- 2) Hold and wait

Starving

instead of sending a wait request processes are put in a job queue

↳ a process can tell in adv. what process it needs (impractical)

↳ a process while requesting for a resource must hold all the resources need by it

3) NO preemption

Resources

impractical as the process might be middle of execution OS and steadily OS takes away the resource

4) Circular wait

every process can take resources in increasing order (in number the resources)

Req. Deadlock Avoidance

	Allocated	max
P ₀	R ₀ R ₁	R ₀ R ₁
P ₁	0 1	7 5
P ₂	2 0	3 2
P ₃	3 0	9 0
P ₄	1 1	2 2
P ₅	0 0	4 3

- Total = $\langle 10, 5 \rangle$
R₀ R₁

P₃ → 2 1

Available → $\langle 3, 2 \rangle$

to check if safe state is true or not we generate a safe sequence

Safe sequence

→ A sequence where the process will run one after another and all the processes will get resources.

Request made by P_i can still be satisfied by current available resources + resources held by previous processes.

If there exist even one safe sequence the process will never go in deadlock.

We need to generate need matrix first.

Basic check

Consider a scenario where P₃ request for R₃ $\langle 1, 0 \rangle$ would OS grant this request?

• Total allocation is not ^{max} that maximum

• resources are available or not

Ex → $\langle 4, 3 \rangle$ are available R₀ R₁.

• OS assumes that resources are allocated and check that do we have safe state after allocation

	Allocated	Max	Need
P ₀	0	7	7
P ₁	2	0	3
P ₂	3	0	9
P ₃	2	4	2
P ₄	0	0	4

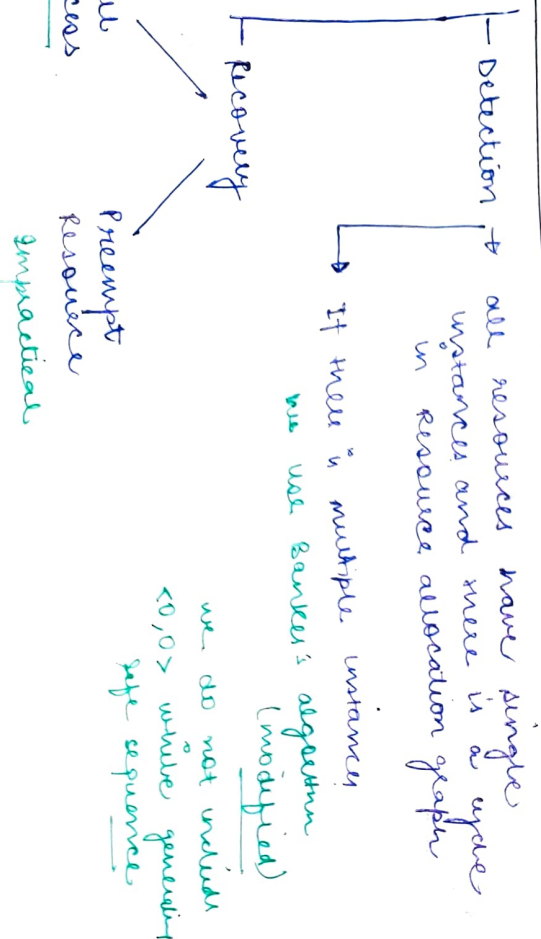
Algorithm

```

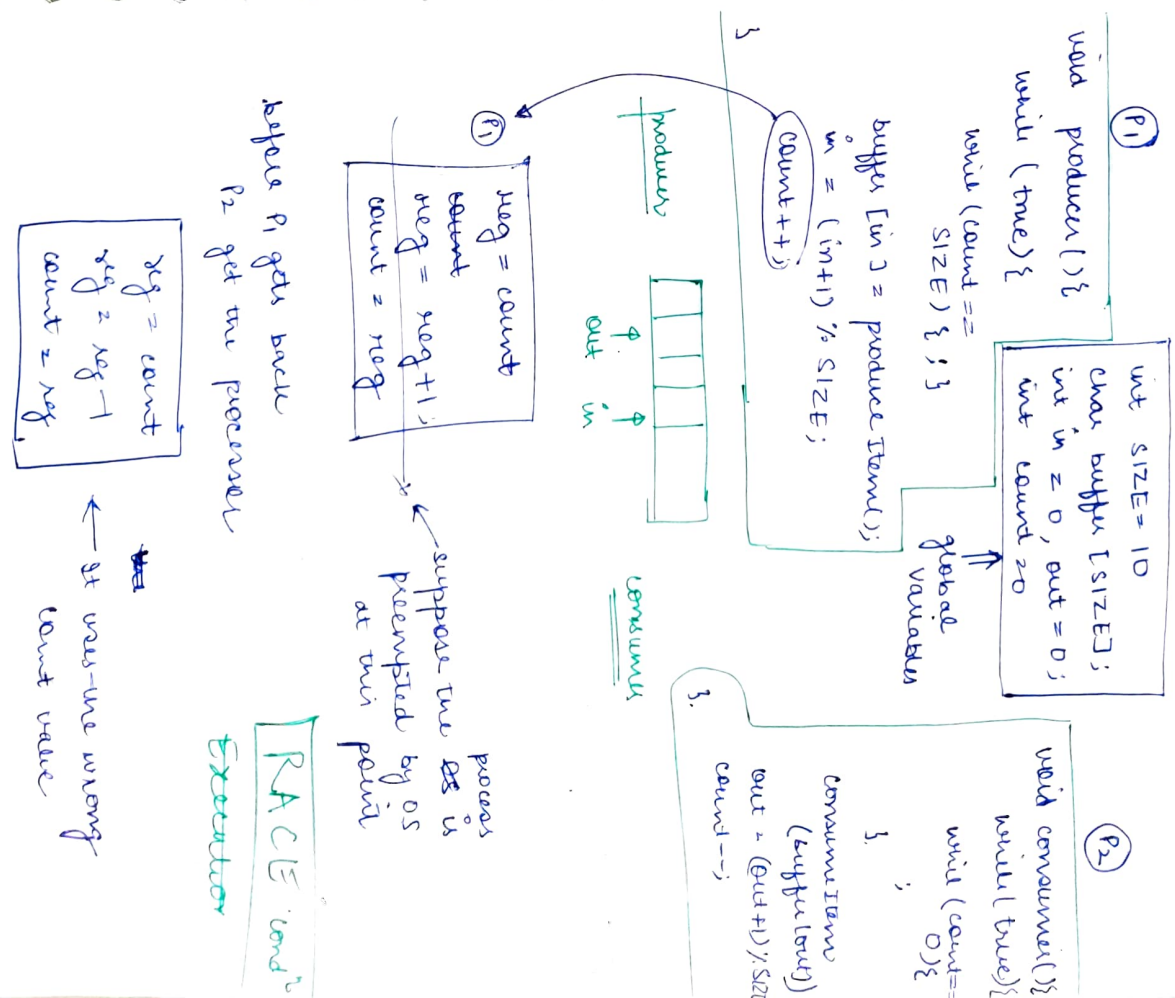
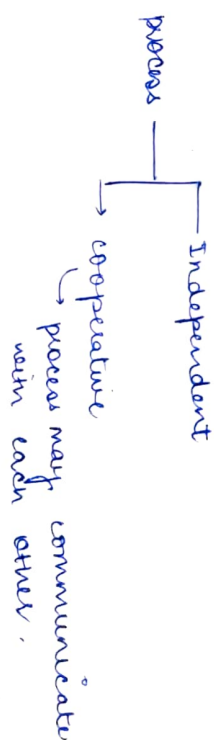
safe-seq = {}
while (All Pi are added) {
  (a) find Pi such that needi ≤ available
  (b) if (no such i exist)
    return false
  else (we found i) {
    available += allocatedi
    Add Pi to the safe-seq
  }
}
return true;

```

Deadlock detection and recovery



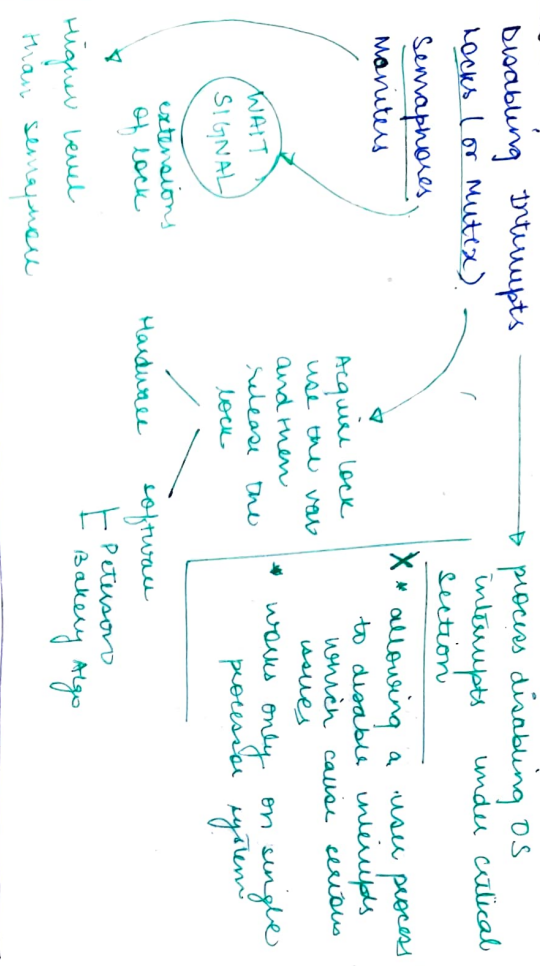
Process synchronization



Critical section → The section within mandatory have to run as a whole. is put in this section

part of program which tries to access shared resources

Synchronization mechanism



Lock for synchronization

```
int amount = 100;
bool lock = false;
```

```
void deposit (int x) {
    while (lock == true) {
        // waiting
    }
    lock = true;
    amount = amount + x;
    lock = false;
}
```

```
void withdraw (x) {
    while (lock == true) {
        // waiting
    }
    lock = true;
    amount = amount - x;
    lock = false;
}
```

```
void deposit (x) {
    while (test_and_set (lock)) {}
    amount = amount + x;
    lock = false;
}
```

```
void withdraw (x) {
    while (test_and_set (lock)) {}
    amount = amount - x;
    lock = false;
}
```

Atomic

```
bool test_and_set (bool *ptr) {
    bool old = *ptr;
    *ptr = true;
    return old;
}
```

Semaphore

```
struct sem_t {
    int count;
    queue q;
```

wait() → decreases count (resources used)

signal() → increases count (resources available)



In the queue the semaphore store the PCBs of process

if value of count is negative, it will assign the process to the queue.

if value of count is negative, it will show the number of process waiting in the queue

- As any process approach and try to use the resource semaphore reduces the count and when it is negative
- if negative, it shows the process into in queue
- if positive already the process

void wait () {

s.count--;

if (s.count < 0) {

① Add the caller process to q;

② sleep(P);

}

}

void signal () {

s.count++;

if (s.count < 0) {

① Remove the process P from q;

② wakeup(P);

}

Binary Semaphore

the count variable is of type bool either 0 or 1.

struct binsem {

bool val;

queue q;

}

void wait () {

if (s.val == 1)

s.val = 0;

else {

1) put this process in queue

2) sleep(P);

}

}

binsem(true, empty);

void signal () {

if (q is empty)

s.val = 1;

else {

1) pick a process from q

2) wakeup(P);

}

}

Monitors

→ make a class which contains shared variables and functions which will use that are synchronized.

Memory Management in OS

Memory Hierarchy

CPU

Cache

→ fast accessible
→ low capacity

Main memory

→ wider of both

Secondary memory

→ lowest highest accessible time

→ fast accessible
→ storing cost low
→ high capacity

Address Binding

→ Binary of a program contains relocatable addresses.

Executable binary

001	-	-	-
002	-	-	-
003	-	-	-
004	-	-	-
005	-	-	-
goto	002		

• To run the binary file it must be brought in main memory and address in main memory is called physical add.

Address binding → mapping of relocatable address to physical add.

Compile time

Load time

Run-time

Compiler knows where the program is to be loaded in phys. mem.

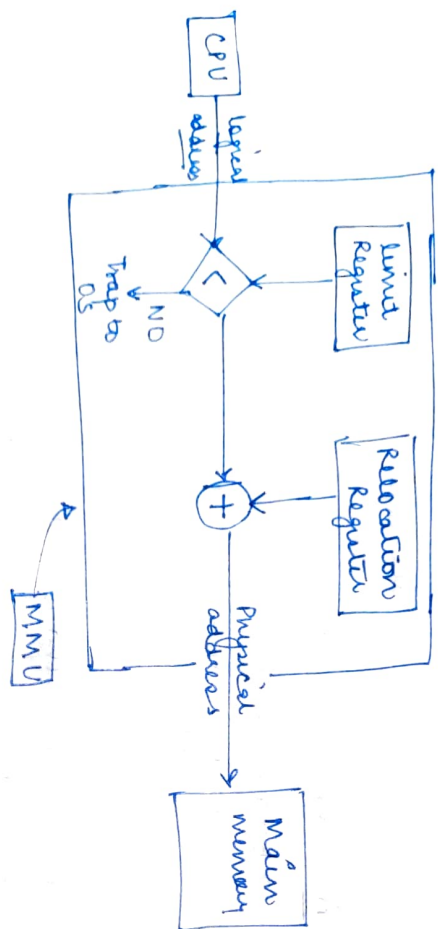
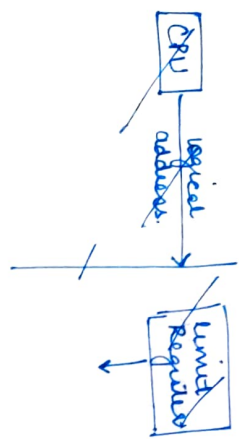
Load time: binary ex. files are loaded into mem, it is wrapped against phys. main memory

Run-time: process can be moved during run-time. address generated by CPU are not phys. address; instead logical address is generated.

MMU → convert LA to PA

Run-time Binding

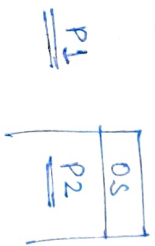
- CPU generates logical address.
- happens through hardware
- Hardware → MMU → memory management unit



we can implement this using software too, but for every context switch OS has to run process (P) that convert logical to physical. And, doing this for every instruction is a costly step.

Evolution of memory management.

Single tasking



we want memory to hold more than 1 process at a time

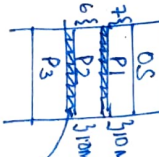
2) Multi-tasking system



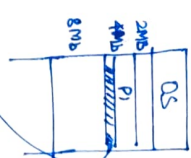
memory allocation

Static

Equal size



Unequal size



memory lens is huge. External Fragmentation

memory is sufficient to hold a process but divided in chunks.

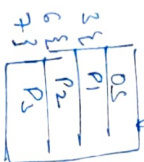
Internal Fragmentation

defragmentation
reorganizing the data stored

Dynamic

(i) contiguous allocation process

the next av. memory



Problems

suppose, P2 leaves the memory now the space left cannot be utilized by another process.

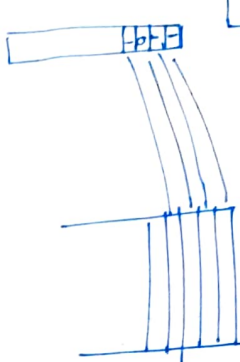
(ii) Non-contiguous

- (a) paging
- (b) segmentation
- (c) paging with segmentation.

Dynamic Partitioning

the holes created by the leaving of processes has to be managed efficiently

1) Bitmap



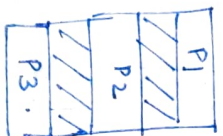
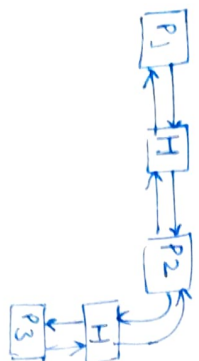
Requires a lot of space



memory is divided into units and then mapped with the corresponding block in bitmap

2) Linked list

where block is represented by a block of LL



- 1) First fit
- 2) Best fit
- 3) Next fit
- 4) Worst fit



travel whole list to get best fit? waste small holes

search for the first hole we

not can accommodate the process.

Next fit → begin from the place we stopped previously

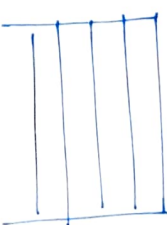
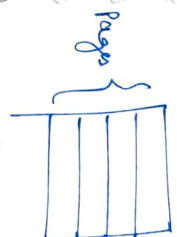
Worst fit → always allocate the biggest size left.

First fit → Best

Paging in memory management

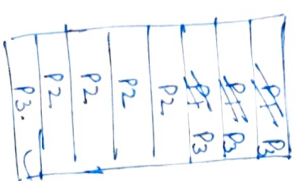
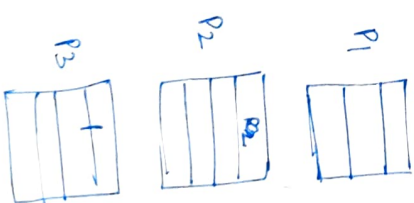
Internally program is stored in fragments

Main memory →

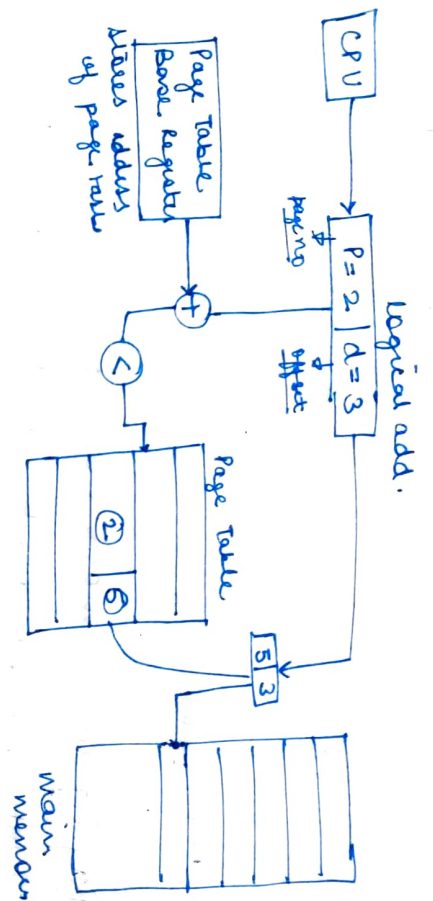


Frames of equal size

→ when you load a process in memory, its individual pages may be loaded at cont. locations or non



not stored contiguously



Page table stores the mapping of logical address pages to mem. frames.

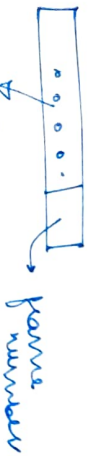
$P=2, d=3$

accesses page 2 and within page 2 go to offset 3

page size
 → small
 → very big
 → page table
 → small
 → very big
 → page table
 → small
 → very big
 → page table

general overhead size → 4Kb

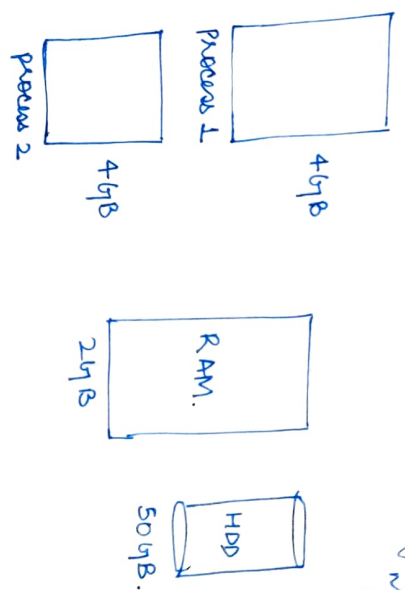
All the pages may not be present in main mem.
 → **Page fault**
 To keep track we have a bit called valid invalid bit which keeps track of the pages (stored or swapped out)
 → That page has to be called from hard disk



modified bit (tells if this page is modified in RAM)

Virtual Memory

→ some pages of process may be present in memory & some may not.



Only load the required part of the process in the main memory and if any required page is not found (page fault), then the page is brought into memory from hard disk.

Page fault → costly operation

→ media switching reading data from HD

can increase the access time by a huge ratio.

pure demand paging →

load the page from HDD when needed, and only load the required part of process in memory.

locality of reference

→ load some nearby pages also.

TLB → Translational lookaside buffer

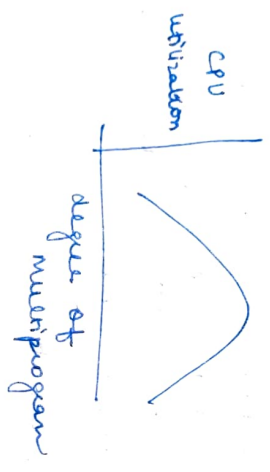
↓
CPU generates logical address.
→ LA has to be converted into PA
i.e. we have to look this PA into page table & this lookup is very frequent and we can optimize it using TLB.

Demand Paging

↳ only req. waking lot of process in memory

thrashing

↳ meaning is the condition when CPU utilization goes down due to high level of ~~program~~ multiprogramming / page faults



Page Replacement Algo

move existing page from main memory to accommodate the demanded page.

- 1) FIFO (suffer from Belady's anomaly).
- 2) optimal
- 3) least recently used.

you remove the page which is going to use later in future (preferable)

replace the earliest used page in time

if the no. of frames is increased, no. of page faults inc.

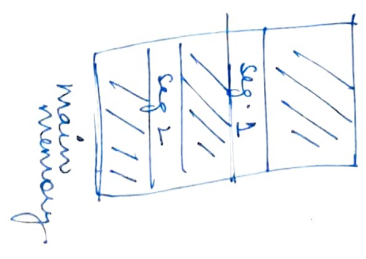
segmentation in OS

→ brought the related items together
we divided the process according to user's view, not necessary equal size

logical view

segment 1
segment 2
segment 3
segment 4

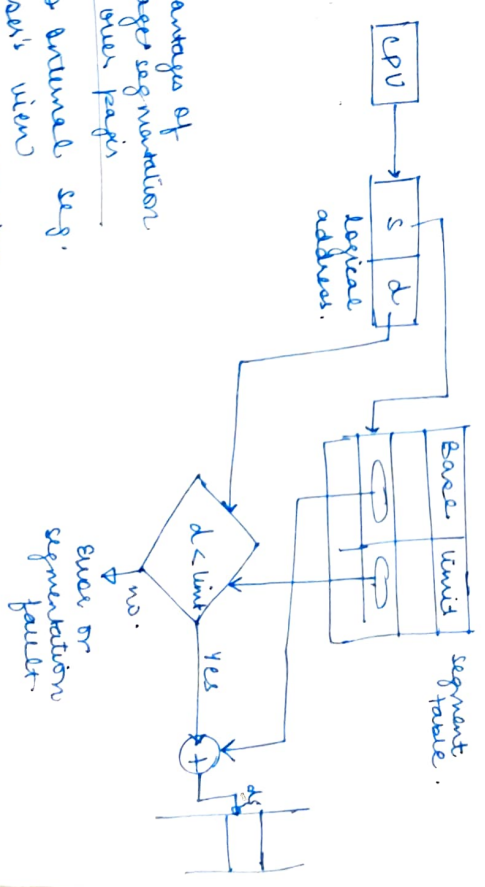
Base	limit	A/P
1600	4000	1
1000	1500	1
///	///	0
///	///	0



initial memory → all segment used, not to be segmentation
segmentation present in main memory.

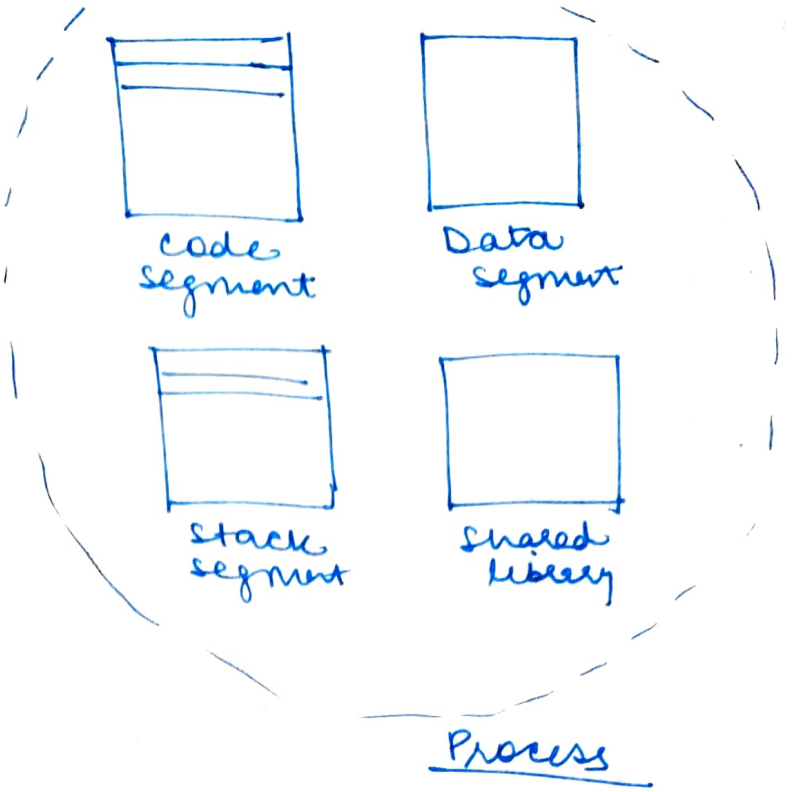
Advantages of page segmentation over page

- No external seg.
- user's view
- Segment table size is smaller than page table size



Disadvantages → External Fragmentation

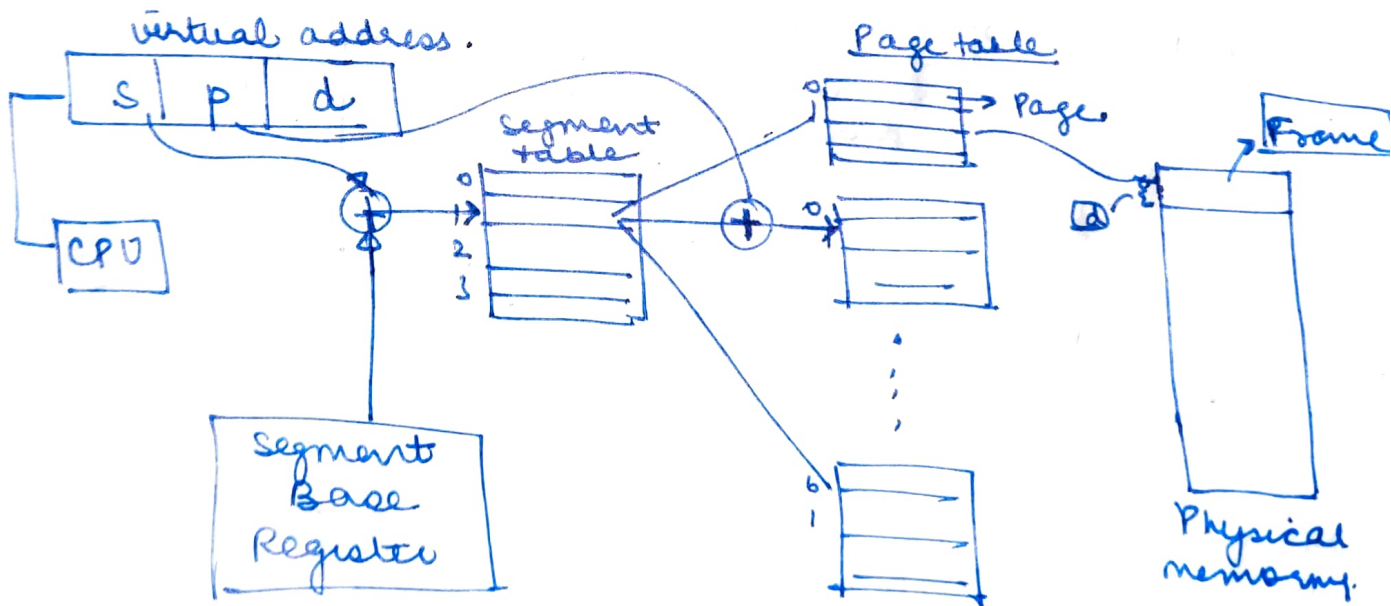
Segmentation with paging



→ Divided into segments and individual segments have pages.

logical address

→ segment number
→ page number
→ offset



disadvantage → two lookups